

DEWEY B. LARSON



BASIC
PROPERTIES
OF MATTER

BASIC PROPERTIES OF MATTER

OTHER BOOKS BY DEWEY B. LARSON

Physical Science

The Structure of the Physical Universe
The Case Against the Nuclear Atom
Beyond Newton
New Light on Space and Time
Quasars and Pulsars
Nothing But Motion
The Universe of Motion
The Neglected Facts of Science
The Liquid State Papers
An Introduction to the Reciprocal System

Economic Science

The Road to Full Employment
The Road to Permanent Prosperity

Metaphysics

Beyond Space and Time

BASIC PROPERTIES OF MATTER

Volume II
of a revised and enlarged edition of
THE STRUCTURE OF THE PHYSICAL
UNIVERSE

By
DEWEY B. LARSON

North Pacific Publishers
Portland, Oregon

Copyright © 1988, 2022
by Ronald & Linda Reilly

All rights reserved

SECOND EDITION
2022

Library of Congress Catalog Card No. 79-88078
ISBN 9798357789266

Contents

	Preface	
1	Solid Cohesion	1
2	Inter-atomic Distances	17
3	Distances in Compounds	35
4	Compressibility	53
5	Heat	77
6	Specific Heat Patterns	91
7	Temperature Relations	111
8	Thermal Expansion	129
9	Electric Currents	141
10	Electrical Resistance	155
11	Thermoelectric Properties	169
12	Scalar Motion	185
13	Electric Charges	201
14	The Basic Forces	219
15	Electrical Storage	233
16	Induction of Charges	251
17	Ionization	265
18	The Retreat from Reality	281
19	Magnetostatics	299
20	Magnetic Quantities and Units	313
21	Electromagnetism	329
22	Magnetic Materials	345
23	Charges in Motion	361
24	Isotopes	377
25	Radioactivity	387
26	Atom Building	401
27	Mass and Energy	415
	References	431

Preface to the Second Printing

More than three decades have passed since the first edition of *Basic Properties of Matter* was released in a limited printing of 200 exemplars. A few words may be in order as to how that printing came about.

At a conference of the International Society of Unified Science held in New York City in the summer of 1986, attended by Dewey Larson and his wife Dorothy, members of the society raised with him the question of the missing volume II of the revised and enlarged edition of the *Structure of the Physical Universe*, originally published in 1959. The first volume of the revised edition, *Nothing but Motion*, had been published in 1979. The third volume, *The Universe of Motion*, followed five years later. In view of the author's advanced age, the members were unanimous in encouraging him to complete the missing intermediary volume without further delay. Larson responded that the work existed in the form of a manuscript, the first eleven chapters of which had been released as a pre-print in 1980, but that the entire work needed a thorough revision in the light of the latest theoretical developments, particularly those involving his recently-completed investigation of the properties of scalar motion, the results of which had been published in a slim volume entitled *The Neglected Facts of Science* in 1982. He was not prepared to undertake the task, however, unless there was a realistic prospect of publication. In response, members of the Society suggested a limited, provisional edition, using the then-new technology of desktop publishing, to be followed, when sufficient funds became available, by a standard book consistent in form and layout with the rest of the series. After some discussion, Larson agreed. I volunteered to take care of the typesetting, Eden Muir to prepare the illustrations, Phillip Porter to handle the printing and binding and Rainer Huck, then Treasurer of the Society, to be in charge of the financing. To facilitate the project, I loaned a Macintosh computer and a printer to the author. Despite his initial reluctance to adopt the

new technology at his advanced age, a few lessons from a family friend were sufficient for him acquire the needed proficiency and before long diskettes began to arrive in the mail with individual chapters of *Basic Properties of Matter*, the name chosen by Larson for the second volume of the series.

The provisional edition of volume II was completed by December 1987 and printed a few months later in 200 numbered exemplars. The result, unfortunately, was not in line with the author's expectations, in view of the poor quality of the printing and binding, and numerous typographical errors. He sent me a list of corrections to be incorporated in the definitive edition that had been promised, but for a variety of reasons, no such edition was published during the author's lifetime or during the subsequent decades.

It is only recently, after my retirement, that I have had the opportunity to prepare for publication a corrected edition of *Basic Properties of Matter*, volume II of the revised and expanded edition of *The Structure of the Physical Universe*. In the process, I have thoroughly reviewed the author's final manuscript, in order to ensure that it is free of the sorts of errors that marred the 1988 provisional edition, and thus provide a reliable tool for researchers.

The book constitutes the final development of the Reciprocal System of theory, based on a lifetime of systematic research and incorporating the important new perspectives on scalar motion and field theory that Larson reached in the 1980s. Despite inevitable physical frailty, he was at the time at the pinnacle of his intellectual powers, as all those who knew him personally can attest.

March 2021

Jan Sammer

Preface

This volume is the second in a series in which I am undertaking to develop the consequences that necessarily follow if it is postulated that the physical universe is composed entirely of motion. The characteristics of the basic motion were defined in *Nothing But Motion*, the first volume of the series, in the form of seven assumptions as to the nature and interrelation of space and time. In the subsequent development, the necessary consequences of these assumptions have been derived by logical and mathematical processes without the introduction of any supplementary or subsidiary assumptions, and without introducing anything from experience. Coincidentally with this theoretical development, it has been shown that the conclusions thus reached are consistent with the relevant data from observation and experiment, wherever a comparison can be made. This justifies the assertion that, to the extent to which the development has been carried, the theoretical results constitute a true and accurate picture of the actual physical universe.

In a theoretical development of this nature, starting from a postulate as to the fundamental nature of the universe, the first results of the deductive process necessarily take the form of conclusions of a basic character: the structure of matter, the nature of electromagnetic radiation, etc. Inasmuch as these are items that cannot be apprehended directly, it has been possible for previous investigators to formulate theories of an ad hoc nature in each individual field to fit the limited, and mainly indirect, information that is available. The best that a *correct* theory can do in any *one* of these individual areas is to arrive at results that *also* agree with the available empirical information. It is not possible, therefore, to grasp the full significance of the new development unless it is recognized that the new theoretical system, the Reciprocal System, as we call it, is one of *general* application, one that reaches all of its conclusions *all* physical fields by deduction from the *same* set of basic premises.

Experience has indicated that it is difficult for most individuals to get a broad enough view of the fundamentals of the many different branches of physical science for a full appreciation of the unitary character of

this new system. However, as the deductive development is continued, it gradually extends down into the more familiar areas, where the empirical information is more readily available, and less subject to arbitrary adjustment or interpretation to fit the prevailing theories. Thus the farther the development of this new general physical theory is carried, the more evident its validity becomes. This is particularly true where, as in the subject matter treated in this present volume, the theoretical deductions provide both explanations and numerical values in areas where neither is available from conventional sources.

There has been an interval of eight years between the publication of Volume I and the first complete edition of this second volume in the series. Inasmuch as the investigation whose results are here being reported is an ongoing activity, a great deal of new information has been accumulated in the meantime. Some of this extends or clarifies portions of the subject matter of the first volume, and since the new findings have been taken into account in dealing with the topics covered in this volume, it has been necessary to discuss the relevant aspects of these findings in this volume, even though some of them may seem out of place. If, and when, a revision of the first volume is undertaken, this material will be transferred to Volume I.

The first 11 chapters of this volume were published in the form of reproductions of the manuscript pages in 1980. Publication of the first complete edition has been made possible through the efforts of a group of members of the International Society of Unified Science, including Rainer Huck, who handled the financing, Phil Porter, who arranged for the printing, Eden Muir, who prepared the illustrations, and Jan Sammer, who was in charge of the project.

December 1987

D. B. Larson

CHAPTER 1

Solid Cohesion

The consequences of the reversal of direction (in the context of a fixed reference system) that takes place at unit distance were explained in a general way in chapter 8 of Volume I. As brought out there, the most significant of these consequences is that establishment of an equilibrium between gravitation and the progression of the natural reference system becomes possible.

There is a location *outside* unit distance where the magnitudes of these two motions are equal: the distance that we are calling the gravitational limit. But this point of equality is not a point of equilibrium. On the contrary, it is a point of instability. If there is even a slight unbalance of forces one way or the other, the resulting motion accentuates the unbalance. A small inward movement, for instance, strengthens the inward force of gravitation, and thereby causes still further movement in the same direction. Similarly, if a small outward movement occurs, this weakens the gravitational force and causes further outward movement. Thus, even though the inward and outward motions are equal at the gravitational limit, this is actually nothing but a point of demarcation between inward and outward motion. It is not a point of equilibrium.

In the region *inside* unit distance, on the contrary, the effect of any change in position opposes the unbalanced forces that produced the change. If there is an excess gravitational force, an outward motion occurs which weakens gravitation and eliminates the unbalance. If the gravitational force is not adequate to maintain a balance, an inward motion takes place. This increases the gravitational effect and restores the equilibrium. Unless there is some intervention by external forces, atoms move gravitationally until they eventually come within unit distance of other atoms. Equilibrium is then

established at positions within this inside region: the time region, as we have called it.

The condition in which a number of atoms occupy equilibrium positions of this kind in an aggregate is known as the *solid state* of matter. The distance between such positions is the *inter-atomic distance*, a distinctive feature of each particular material substance that we will examine in detail in the following chapter. Displacement of the equilibrium in either direction can be accomplished only by the application of a force of some kind, and a solid structure resists either an inward force, a *compression*, or an outward force, a *tension*. To the extent that resistance to tension operates to prevent separation of the atoms of a solid it is commonly known as the force of *cohesion*.

The conclusions with respect to the nature and origin of atomic cohesion that have been reached in this work replace a familiar theory, based on altogether different premises. This previously accepted hypothesis, the *electrical theory of matter*, has already had some consideration in the preceding volume, but since the new explanation of the nature of the cohesive force is basic to the present development, some more extensive comparisons of the two conflicting viewpoints will be in order before we proceed to develop the new theoretical structure in greater detail.

The electrical, or electronic, theory postulates that the atoms of solid matter are electrically charged, and that their cohesion is due to the attraction between unlike charges. The principal support for the theory comes from the behavior of ionic compounds in solution. A certain proportion of the molecules of such compounds split up, or *dissociate*, into oppositely charged components which are then called *ions*. The presence of the charges can be explained in either of two ways: (1) the charges were present, but undetectable, in the undissolved material, or (2) they were created in the solution process. The adherents of the electrical theory base it on explanation (1). At the time this explanation was originally formulated, electric charges were thought to be relatively permanent entities, and the conclusion with respect to their role in the solution process was therefore quite in keeping with contemporary scientific thought. In the meantime, however, it has been found that electric charges are easily created and easily destroyed, and

are no more than a transient feature of matter. This cuts the ground from under the main support of the electrical theory, but the theory has persisted because of the lack of any available alternative.

Obviously *some* kind of a force must hold the solid aggregate together. Outside of the forces known to result directly from observable motion, there are only three kinds of force of which there has heretofore been any definite observational knowledge: gravitational, electric, and magnetic. The so-called “forces” which play various roles in present-day atomic physics are purely hypothetical. Of the three known forces, the only one that appears to be strong enough to account for the cohesion of solids is the electric force. The general tendency in scientific circles has therefore been to take the stand that cohesion *must* result from the operation of electrical forces, notwithstanding the lack of any corroboration of the conclusions reached on the basis of the solution process, and the existence of strong evidence against the validity of those conclusions.

One of the serious objections to this electrical theory of cohesion is that it is not actually *a* theory, but a patchwork collection of theories. A number of different explanations are advanced for what is, to all appearances, the same problem. In its basic form, the theory is applicable only to a restricted class of substances, the so-called “ionic” compounds. But the great majority of compounds are “non-ionic.” Where the hypothetical ions are clearly non-existent, an electrical force between ions cannot be called upon to explain the cohesion, so, as one of the general chemistry texts on the author’s shelves puts it, “A different theory was required to account for the formation of these compounds.” But this “different theory,” based on the weird concept of electrons “shared” by the interacting atoms, is still not adequate to deal with all of the non-ionic compounds, and a variety of additional explanations are called upon to fill the gaps.

In current chemical parlance the necessity of admitting that each of these different explanations is actually another theory of cohesion is avoided by calling them different types of “bonds” between the atoms. The hypothetical bonds are then described in terms of interaction of electrons, so that the theories are united in language, even though widely divergent in content. As noted in Chapter 19, Vol. I, a half

dozen or so different types of bonds have been postulated, together with “hybrid” bonds which combine features of the general types.

Even with all of this latitude for additional assumptions and hypotheses, some substances, notably the metals, cannot be accommodated within the theory by any expedient thus far devised. The metals admittedly do not contain *oppositely* charged components, if they contain any charged components at all, yet they are subject to cohesive forces that are indistinguishable from those of the ionic compounds. As one prominent physicist, V. F. Weisskopf, found it necessary to admit in the course of a lecture, “I must warn you I do not understand why metals hold together.” Weisskopf points out that scientists cannot even agree as to the manner in which the theory should be applied. Physicists give us one answer, he says, chemists another, but “neither of these answers is adequate to explain what a chemical bond is.”¹

This is a significant point. The fact that the cohesion of metals is clearly due to something other than the attraction between unlike charges logically leads to a rather strong presumption that atomic cohesion in general is non-electrical. As long as some non-electrical explanation of the cohesion of metals has to be found, it is reasonable to expect that this explanation will be found applicable to other substances as well. Experience in dealing with cohesion of metals thus definitely foreshadows the kind of conclusions that have been reached in the development of the Reciprocal System of theory.

It should also be noted that the electrical theory is wholly ad hoc. Aside from what little support it can derive from extrapolation to the solid state of the conditions existing in solutions, there is no independent confirmation of *any* of the principal assumptions of the theory. No observational indication of the existence of electrical charges in ordinary matter can be detected, even in the most strongly ionic compounds. The existence of electrons as *constituents* of atoms is purely hypothetical. The assumption that the reluctance of the inert gases to enter into chemical compounds is an indication that their structure is a particularly stable one is wholly gratuitous. And even the originators of the idea of “sharing” electrons make no attempt to provide any meaningful explanation of what this means, or how it could be accomplished, if there actually were any electrons in the

atomic structure. These are the assumptions on which the theory is based, and they are entirely without empirical support. Nor is there any solid basis for what little theoretical foundation the theory may claim, inasmuch as its theoretical ties are to the nuclear theory of atomic structure, which is itself entirely ad hoc.

But these points, serious as they are, can only be regarded as supplementary evidence, as there is one fatal weakness of the electrical theory that would demolish it even if nothing else of an adverse nature were known. This is our knowledge of the behavior of positive and negative electric charges when they are brought into close proximity. Such charges do not establish an equilibrium of the kind postulated in the theory; they destroy each other. There is no evidence which would indicate that the result of such contact is any different in a solid aggregate, nor is there even any plausible theory as to *why* any different outcome could be expected, or *how* it could be accomplished.

It is worth noting in this connection that while current physical theory portrays positive and negative charges as existing in a state of congenial companionship in the nuclear theory of the atom and in the electrical theory of matter, it turns around and gives us explanations of the behavior of antimatter in which these charges display the same violent antagonism that they demonstrate to actual observation. This is the kind of inconsistency that inevitably results when recalcitrant problems are “solved” by ad hoc assumptions that involve departures from established physical laws and principles.

In the context of the present situation in which the electrical theory is challenged by a new development, all of these deficiencies and contradictions that are inherent in the electrical theory become very significant. But the positive evidence in favor of the new theory is even more conclusive than the negative evidence against its predecessor. First, and probably the most important, is the fact that we are not replacing the electrical theory of matter with another “theory of matter.” The Reciprocal System is a complete general theory of the physical universe. It contains no hypotheses other than those relating to the nature of space and time, and it produces an explanation of the cohesion of solids in the same way that it derives logical and consistent explanations of other physical phenomena: simply by

developing the consequences of the basic postulates. We therefore do not have to call upon any additional force of a hypothetical nature to account for the cohesion. The two forces that determine the course of events in the region outside unit distance also account for the existence of the inter-atomic equilibrium inside this distance.

It is significant that the new theory identifies *both* of these forces. One of the major defects of the electrical theory of cohesion is that it provides only one force, the hypothetical electrical force of attraction, whereas two forces are required to explain the observed situation. Originally it was assumed that the atoms are impenetrable, and that the electrical forces merely hold them in contact. Present-day knowledge of compressibility and other properties of solids has demolished this hypothesis, and it is now evident that there must be what Karl Darrow called an “antagonist,” in the statement quoted in Volume I, to counter the attractive force, whatever it may be, and produce an equilibrium. Physicists have heretofore been unable to find any such force, but the development of the Reciprocal System has now revealed the existence of a powerful and omnipresent force hitherto unknown to science. Here is the missing ingredient in the physical situation, the force that not only explains the cohesion of solid matter, but, as we saw in Volume I, supplies the answers to such seemingly far removed problems as the structure of star clusters and the recession of the galaxies.

One point that should be specifically noted is that it is this hitherto unknown force, the force due to the progression of the natural reference system, that holds the solid aggregate together, not gravitation, which acts in the opposite direction in the time region. The prevailing opinion that the force of gravitation is too weak to account for the cohesion is therefore irrelevant, whether it is correct or not.

Inasmuch as the new theoretical system applies the same general principles to an understanding of all of the inter-atomic and inter-molecular equilibria, it explains the cohesion of *all* substances by the *same* physical mechanism. It is no longer necessary to have one theory for ionic substances, several more for those that are non-ionic, and to leave the metals out in the cold without any applicable theory. The theoretical findings with respect to the nature of chemical

combinations and the structure of molecules that were outlined in the preceding volume have made a major contribution to this simplification of the cohesion picture, as they have eliminated the need for different kinds of cohesive forces, or “bonds.” All that is now required of a theory of cohesion is that it supply an explanation of the inter-atomic equilibrium, and this is provided, for all solid substances under all conditions, by balancing the outward motion (force) of gravitation against the inward motion (force) of the progression of the natural reference system. Because of the asymmetry of the rotational patterns of the atoms of many elements, and the consequent anisotropy of the force distributions, the equilibrium locations vary not only between substances, but also between different orientations of the same substance. Such variations, however, affect only the magnitudes of the various properties of the atoms. The essential character of the inter-atomic equilibrium is always the same.

As indicated in the original discussion of gravitation, even though the various aggregates of matter do not actually exert gravitational forces on each other, the observable results of their gravitational motions are identical with those that would be produced if such forces did exist. The same is true of the results of the progression of the natural reference system. There is a considerable element of convenience in expressing these results in terms of force, on an “as if” basis, and this practice has already been followed to some extent in the previous volume. Now that we are ready to begin a quantitative evaluation of the inter-atomic relations, however, it is desirable to make it clear that the force concept is being used only for convenience. Although the quantitative discussion that follows, like the earlier qualitative discussion, will be carried on in terms of forces, what we will actually be dealing with are the inward and outward motions of each individual atom.

While the items that have been mentioned add up to a very impressive case in favor of the new theory of cohesion, the strongest confirmation of its validity comes from its ability to *locate* the point of equilibrium; that is to give us specific values of the inter-atomic distances. As will be demonstrated in Chapter 2, we are already able, by means of the newly established relations, to calculate the possible values of the inter-atomic distance for most of the simpler

substances, and there do not appear to be any serious obstacles in the way of extending the calculations to more complex substances whenever the necessary time and effort can be applied to the task. Furthermore, this ability to determine the location of the point of equilibrium is not limited to the simple situation where only the two basic forces are involved. Chapters 4 and 5 will show that the same general principles can also be applied to an evaluation of the changes in the equilibrium distance that result from the application of heat or pressure to the solid aggregate.

Although, as stated in Volume I, the true magnitude of a unit of space is the same everywhere, the *effective* magnitude of a spatial unit in the time region is reduced by the inter-regional ratio. It is convenient to regard this reduced value, $\frac{1}{156.44}$ of the natural unit, as the *time region unit of space*. The effective portion of a time region phenomenon may extend into one or more additional units, in which case the measured distance will exceed the time region unit, or the original single unit may not be fully effective, in which case the measured distance will be less than the time region unit. Thus the inter-atomic equilibrium may be reached either inside or outside the time region unit of distance, depending on where the outward rotational forces reach equality with the inward force of the progression of the natural reference system. Extension of the inter-atomic distance beyond one time region unit does not take the equilibrium system out of the time region, as the boundary of that region is at one full-sized natural unit of distance, not at one time region unit. So far as the inter-atomic force equilibrium is concerned, therefore, the time region unit of distance does not represent any kind of a critical magnitude.

As we saw in our examination of the composition of the magnetic neutral groups, however, the natural unit as it exists in the time region (the time region unit) is a critical magnitude from the orientation standpoint. An explanation of this difference can be derived from a consideration of the difference in the inherent nature of the two phenomena. Where the inter-atomic distance is less than one time region unit, the rotational forces are acting against the inward force of the progression of the reference system during only a portion of the unit progression. Similarly, where the inter-atomic distance is greater than one time region unit, the unit inward force is acting

against only a portion of the greater-than-unit outward rotational forces. The variations in distance thus reflect differences in the *magnitudes* of the rotational forces. But the orientation effect has no magnitude. It either exists, or does not exist. As we have noted in the previous discussion, particularly in connection with the structure of the benzene molecule, this effect, if it exists, is the same regardless of whether it acts at short range or at long range. The essential requirement that it must meet is that it must be *continuously* effective. Otherwise, the orientation is destroyed during the off period. Where the rotational forces extend beyond one time region unit, so that the unit orientation effect is coincident with only a portion of the total rotational forces, the orienting effect is not continuous, and no orientation takes place.

In this chapter we are dealing mainly with what we are calling “rotational forces.” These are, of course, the same “as if” forces due to the scalar aspect of the atomic rotation that were called “gravitational” in some other contexts, the choice of language depending on whether it is the origin or the effect of the force that is being emphasized in the discussion. For a quantitative evaluation of the rotational forces we may use the general force equation, providing that we replace the usual terms of the equation with the appropriate time region terms. As explained in introducing the concept of the time region in Chapter 8 of Vol. I, equivalent space $\frac{1}{t}$ replaces space in the time region, and velocity is therefore $\frac{1}{t^2}$. Energy, the one-dimensional equivalent of mass, which takes the place of mass in the time region expression of the force equation, because the three rotations of the atom act separately, rather than jointly, in this region, is the reciprocal of this expression, or t^2 . Acceleration is velocity divided by time: $\frac{1}{t^3}$. The time region equivalent of the equation $F = ma$ is therefore $F = Ea = t^2 \times \frac{1}{t^3} = \frac{1}{t}$ in each dimension.

At this point we will need to take note of the nature of the increments of speed displacement in the time region. In the outside region additions to the displacement proceed by units: first one unit, then another similar unit, yet another, and so on, the total up to any specific point being n units. There is no term with the value n . This value appears only as a total. The additions in the time region follow a different mathematical pattern, because in this case only one of

the components of motion progresses, the other remaining fixed at the unit value. Here the displacement is $1/x$, and the sequence is $1/1$, $1/2$, $1/3 \dots 1/n$. The quantity $1/n$ is the final term, not the total. To obtain the total that corresponds to n in the outside region it is necessary to integrate the quantity $1/x$ from $x = 1$ to $x = n$. The result is $\ln n$, the natural logarithm of n .

Many readers of the first edition have asked why this total should be an integral rather than a summation. The answer is that we are dealing with a continuous quantity. As pointed out in the introductory chapters of the preceding volume, the motion of which the universe is constructed does not proceed in a succession of jumps. Even though it exists only in units, it is a continuous progression. A unit of this motion is a specific portion of this continuity. A series of units is a more extended segment of that continuity, and its magnitude is an integral. In dealing with the basic individual units of motion in the outside region it is possible to use the summation process, but only because in this case the sum is the same as the integral. To get the total of the $1/x$ series we must integrate.

To evaluate the rotational force we integrate the quantity $1/t$ from unity, the physical datum or zero level, to t :

$$F_r = \int_1^t \frac{1}{t} dt = \ln t \quad (1-1)$$

If the quantity $\ln t$ is below unity in any dimension there is no effective outward force in that dimension, but the natural logarithm exceeds unity for all values of x above 2, and the atoms of all elements have a rotational displacement of 2 (equivalent to $t=3$) or more in at least one dimension. Consequently, all have effective rotational forces.

The force computed from equation 1-1 is the inherent rotational force of the individual atom; that is, the one-dimensional force which it exerts against a single unit of force. The force between two (apparently) interacting atoms is

$$F = \ln t_A \ln t_B \quad (1-2)$$

For a two-dimensional magnetic rotation this becomes

$$F = \ln^2 t_A \ln^2 t_B \quad (1-3)$$

As we found in Chapter 12, Vol. I, the equivalent of distance s in the time region is s^2 , and the gravitational force in this region therefore varies inversely as the fourth power of the distance rather than the square. Applying this factor to the expression for the force of the two-dimensional rotation, together with the inter-regional ratio, the ratio of effective to total force derived in the same chapter, we obtain the effective force of the magnetic rotation of the atom:

$$F_m = (0.006392)^4 s^{-4} \ln^2 t_A \ln^2 t_B \quad (1-4)$$

The distance factor does not apply to the force due to the progression of the natural reference system, as this force is omnipresent, and unlike the rotational force, is not altered as the objects to which it is applied change their relative positions. At the point of equilibrium, therefore, the rotational force is equal to the unit force of the progression. Substituting unity for F_m in equation 1-4, and solving for the equilibrium distance, we obtain

$$s_0 = 0.006392 \ln^{1/2} t_A \ln^{1/2} t_B \quad (1-5)$$

The inter-atomic distances for those elements which have no electric rotation, the inert gas series, may be calculated directly from this equation. In the elements, however, $t_A = t_B$ in most cases, and it will be convenient to express the equation in the simplified form:

$$s_0 = 0.006392 \ln t \quad (1-6)$$

The values thus calculated are in the neighborhood of 10^{-8} cm, and for convenience this quantity has been taken as a unit in which to express the inter-atomic and inter-molecular distances. When converted from natural units to this conventional unit, the Angstrom unit, symbol Å, equation 1-6 becomes

$$s_0 = 2.914 \ln t \text{ Å} \quad (1-7)$$

In applying this equation we encounter another of the questions with respect to terminology that inevitably arise in a basically new treatment of any subject. The significance of the quantity t as used in the foregoing discussion and in the equations is obvious from the context—it is the magnitude of the *effective* rotation—but the question is: What shall we call it? The basic quantity with which we are dealing, the rotational speed displacement, does not enter into the equations directly. The

mathematical structure of these equations requires us to enter them with values that include the initial unit which constitutes the natural zero datum. Furthermore, each double vibrational unit rotates independently, and when the rotation extends to a second such unit the increment in the value of t is only one half unit per added unit of displacement. Under these circumstances, where the relation of the term t to the displacement is variable, it seems advisable to give this term a distinctive name, and we will therefore call it the *specific rotation*.

As brought out in the discussion of the general characteristics of the atomic rotation in Chapter 10, Vol. I, the two magnetic displacements may be unequal, and in this event the speed distribution takes the form of a spheroid with the principal rotation effective in two dimensions and the subordinate rotation in one. The average effective value of the specific rotation under these conditions is $(t_1^2 t_2^2)^{1/3}$. In this case we are dealing with the properties of a single entity, and the mathematical situation seems clear. But it is not so evident how we should arrive at the effective specific rotation where there is an interaction between two atoms whose individual rotations are different. As matters now stand it appears that the geometric mean of the two specific rotations is the correct quantity, and the values tabulated in Chapters 2 and 3 have been calculated on this basis. It should be noted, however, that this conclusion as to the mathematics of the combination is still somewhat tentative, and if further study shows that it must be modified in some, or all, applications, the calculated values will be subject to corresponding modifications. Any changes will be small in most cases, but they will be substantial where there is a large difference between the two components. The absence of major discrepancies between the calculated and observed distances in combinations of atoms with much different dimensions therefore gives some significant support to the use of the geometric mean pending further theoretical clarification.

The inter-atomic distances of four of the five inert gas elements for which experimental data are available follow the regular pattern. The values calculated for these elements are compared with the experimental distances in Table 1.

TABLE 1: *DISTANCES-INERT GAS ELEMENTS*

Atomic Number	Element	Specific Rotation	Distance	
			Calc.	Obs.
10	Neon	3-3	3.21	3.20
18	Argon	4-3	3.74	3.84
36	Krypton	4-4	4.04	4.02
54	Xenon	4½-4½	4.35	4.41

Helium, which also belongs to the inert gas series, has some special characteristics due to its low rotational displacement, and will be discussed in connection with other elements affected by the same factors. The reason for the appearance of the 4½ value in the xenon rotation will also be explained shortly.

The calculated distances are those which would prevail in the absence of compression and thermal expansion. A few of the experimental data have been extrapolated to this zero base by the investigators, but most of them are the actual observed values at atmospheric pressure and at temperatures which depend on the properties of the substances under examination. These values are not exactly comparable to the calculated distances. In general, however, the expansion and compression up to the temperature and pressure of observation are small. A comparison of the values in the last two columns of Table 1 and the similar tables in chapters 2 and 3 therefore gives a good picture of the extent of agreement between the theoretical figures and the experimental results.

Another point about the distance correlations that needs to be taken into account is that there is a substantial amount of variation in the experimental results. If we were to take the closest of these measured values as the basis for comparison, the correlation would be very much better. One relatively recent determination of the xenon distance, for example, arrives at a value of 4.34, almost identical with the calculated distance. There are also reported values for the argon distance that agree more closely with the theoretical result. However, a general policy of using the closest values would introduce a bias that would tend to make the correlation look more favorable than the situation actually warrants. It has therefore been considered advisable to use empirical data from a recognized selection of preferred values. Except for those values identified by asterisks, all

of the experimental distances shown in the tables are taken from the extensive compilation by Wyckoff.² Of course, the use of these values selected on the basis of indirect criteria introduces a bias in the unfavorable direction, since, if the theoretical results are correct, every experimental error shows up as a discrepancy, but even with this negative bias the agreement between theory and observation is close enough to show that the theoretical determination of the inter-atomic distance is correct in principle, and to demonstrate that, with the exception of a relatively small number of uncertain cases, it is also correct in the detailed application.

Turning now to the elements which have electric as well as magnetic displacement, we note again that the electric rotation is one-dimensional and opposes the magnetic rotation. We may therefore obtain an expression for the effect of the electric rotational force on the magnetically rotating photon by inverting the one-dimensional force term of equation 1-2.

$$F_e = \frac{1}{\ln t'_A \ln t'_B} \quad (1-8)$$

Inasmuch as the electric rotation is not an independent motion of the basic photon, but a rotation of the magnetically rotating structure in the reverse direction, combining the electric rotational force of equation 1-8 with the magnetic rotational force of equation 1-4 modifies the rotational terms (the functions of t) only, and leaves the remainder of equation 1-4 unchanged.

$$F = (0.006392)^4 \frac{\ln^2 t_A \ln^2 t_B}{s^4 \ln t'_A \ln t'_B} \quad (1-9)$$

Here again the effective rotational (outward) and natural reference system progression (inward) forces are necessarily equal at the equilibrium point. Since the force of the progression of the natural reference system is unity, we substitute this value for F in equation 1-9 and solve for s_0 , the equilibrium distance, as before.

$$s_0 = 0.006392 \frac{\ln^{1/2} t_A \ln^{1/2} t_B}{\ln^{1/4} t'_A \ln^{1/4} t'_B} \quad (1-10)$$

Again simplifying for application to the elements, where A is generally equal to B ,

$$s_0 = 0.006392 \frac{\ln t}{\ln^{1/2} t}, \quad (1-11)$$

In Angstrom units this becomes

$$s_0 = 2.914 \frac{\ln t}{\ln^{1/2} t}, \text{ \AA} \quad (1-12)$$

As already noted, when the rotation is extended to a second (double) vibrational unit, to *vibration two*, we may say, each added displacement unit adds only one half unit to the specific rotation. Inasmuch as 8 electric displacement units distributed three-dimensionally bring the rotation to a new zero point, and cause the rotational motion to revert to the translational status, the change to vibration two in the electric dimension *must* take place before the displacement reaches 8. Specific rotation 8 (displacement 7) is therefore followed by $8\frac{1}{2}$, 9, $9\frac{1}{2}$, etc. But the first effective rotational displacement unit is necessarily one-dimensional, and the linear equivalent of the 8-unit limit is 2 units. Thus this first unit has already reached the one-dimensional limit. The succeeding displacement units have the option of continuing on the one-dimensional basis and extending the rotation to vibration two rather than extending it into additional dimensions. The change to vibration two therefore *may* take place immediately after the first displacement unit. In this case specific rotation 2 (displacement 1) is followed by $2\frac{1}{2}$, 3, $3\frac{1}{2}$, etc. The lower value is commonly found where it first becomes possible; that is, displacement 2 normally corresponds to rotation $2\frac{1}{2}$ rather than 3. The next element may take the intermediate value $3\frac{1}{2}$, but beyond this point the higher vibration one value normally prevails.

In the first edition it was indicated that the one or two vibrational displacement units being rotated did not necessarily constitute the entire vibrational component of the basic photon, inasmuch as these one or two units are capable of being rotated independently of the remaining vibrational units, if any. Further consideration now leads to the conclusion that one or two units of a multi-unit photon frequency can, in fact, be set in rotation independently, as previously indicated, and that the original photon may have had an excess of vibrational units, but that in such an event the rotating portion of the photon begins moving inward, whereas the non-rotating portion

continues moving outward by reason of the progression of the natural reference system. The two portions therefore separate, and the rotating portion retains no non-rotating vibrational component.

The general pattern of the magnetic rotational values is the same as that of the electric values. The tendency to substitute specific rotation $2\frac{1}{2}$ for 3 applies to the magnetic as well as to the electric rotation, and in the lower group combinations (both elements and compounds) that follow the regular electropositive pattern the specific magnetic rotations are usually $2\frac{1}{2}$ - $2\frac{1}{2}$ or 3 - $2\frac{1}{2}$, rather than 3 - 3 . But the upper limit for specific magnetic rotation on a vibration one basis is 4 (3 displacement units) instead of 8, as the two-dimensional rotation reaches the upper zero level at 4 displacement units in each dimension. Rotation $4\frac{1}{2}$ therefore follows rotation 4 in the regular sequence, as we saw in the values given for xenon in Table 1. It is possible to reach rotation 5 in one dimension, however, without bringing the magnetic rotation as a whole up to the 5 level, and 5 - 4 or 5 - $4\frac{1}{2}$ rotation occurs in some elements either in lieu of, or in combination with, the $4\frac{1}{2}$ - 4 or $4\frac{1}{2}$ - $4\frac{1}{2}$ rotation.

CHAPTER 2

Inter-atomic Distances

As equation 1-10 indicates, the distance between any two atoms in a solid aggregate is a function of the specific rotations of the atoms. Since each atom is capable of assuming any one of several different relative orientations of its rotational motions, it follows that there are a number of possible specific rotations for each combination of atoms. This number of possible alternatives is still further increased by two additional factors that were discussed earlier. The atom has the option, as we noted in Chapter 10, vol. I, of rotating with the normal magnetic displacement and a positive electric increment, or with the next higher magnetic displacement and a negative electric increment. And in either case, the effective quantity, the specific rotation, may be modified by extension of the motion to a second vibrating unit, as brought out in Chapter 1.

It is possible that each of these many variations of the magnitude of the specific rotation, and the corresponding values of the inter-atomic distances, may actually be realized under appropriate conditions, but in any particular set of circumstances certain combinations of rotations are more probable than the others, and in ordinary practice the number of different values of the distance between the same two atoms is relatively small, except in certain special cases. As matters now stand, therefore, we are able to calculate from theoretical premises a small set of possible inter-atomic distances for each element or compound.

Ultimately it will no doubt be advisable to evaluate the probability relations in detail so that the results of the calculations will be as specific as possible, but it has not been feasible to undertake this full treatment of the probability relationships in this present work. In an investigation of so large a field as the structure of the physical universe there must not only be some selection of the subjects that

are to be covered, but also some decisions as to the extent to which that coverage will be carried. A comprehensive treatment of the probability relations wherever they enter into physical situations could be quite helpful, but the amount of time and effort required to carry out such a project will undoubtedly be enormous, and its contribution to the major objectives of this present undertaking is not sufficient to justify allocating so much of the available resources to it. Similar decisions as to how far to carry the investigation in certain areas have had to be made from time to time throughout the course of the work in order to limit it to a finite size.

It might be well to point out in this connection that it will *never* be possible to calculate a unique inter-atomic distance for every element or combination of elements, even when the probability relations have been definitely established, as in many cases the choice from among the alternatives is not only a matter of relative probability, but also of the *history* of the particular specimen. Where two or more alternative forms are stable within the range of physical conditions under which the empirical examination is being made, the treatment to which the specimen has previously been subjected plays an important part in the determination of the structure.

It does not follow, however, that we are totally precluded from arriving at definite values for the inter-atomic distances. Even though no quantitative evaluation of the relative probabilities of the various alternatives is yet available, the nature of the major factors involved in their determination can be deduced theoretically, and this qualitative information is sufficient in most cases to exclude all but a very few of the total number of possible variations of the specific rotations. Furthermore, there are some series relations by means of which the range of variability can be still further narrowed. These series patterns will be more evident when we examine the distances in compounds in the next chapter, and they will be given more detailed consideration at that point.

The first thing that needs to be emphasized as we begin our analysis of the factors that determine the inter-atomic distance is that we are not dealing with the *sizes* of atoms; what we are undertaking to do is to evaluate the distance between the *equilibrium positions* that the

atoms occupy under specified conditions. In Chapter 1 we examined the general nature of the atomic equilibrium. In this and the following chapter we will see how the various factors involved in the relations between the rotations of the (apparently) interacting atoms affect the point of equilibrium, and we will arrive at values of the inter-atomic distances under static conditions. Then in Chapters 5 and 6 we will develop the quantitative relations that will enable us to determine just what changes take place in these equilibrium distances when external forces in the form of pressure and temperature are applied.

As we have seen in the preceding volume, all atoms and aggregates of matter are subject to two opposing forces of a general nature: gravitation and the progression of the natural reference system. These are the primary forces (or motions) that determine the course of physical events. Outside the gravitational limits of the largest aggregates, the outward motion due to the progression of the natural reference system exceeds the inward motion of gravitation, and these aggregates, the major galaxies, move outward from each other at speeds increasing with distance. Inside the gravitational limits the gravitational motion is the greater, and all atoms and aggregates move inward. Ultimately, if nothing intervenes, this inward motion carries each atom within unit distance of another, and the directional reversal that takes place at the unit boundary then results in the establishment of an equilibrium between the motions of the two atoms. The inter-atomic distance is the distance between the atomic centers in this equilibrium condition. It is not, as currently assumed, an indication of the sizes of the atoms.

The current theory which regards the inter-atomic distance as a measure of "size" is, in many respects, quite similar to the electronic "bond" theory of molecular structure. Like the electronic theory, it is based on an erroneous assumption—in this case, the assumption that the atoms are in contact in the solid state—and like the electronic theory, it fits only a relatively small number of substances in its simple form, so that it is necessary to call upon a profusion of supplementary and subsidiary hypotheses to explain the deviations of the observed distances from what are presumed to be the primary values. As the textbooks point out, even in the metals, which are the simplest structures from the standpoint of the theory, there are many difficult

problems, including the awkward fact that the presumed “size” is variable, depending on the nature of the crystal structure. Some further aspects of this situation will be considered in Chapter 3.

The resemblance between these two erroneous theories is not confined to the lack of adequate foundations and to the nature of the difficulties that they encounter. It also extends to the resolution of these difficulties, as the *same* principles that were derived from the postulates of the Reciprocal System to account for the formation of molecules of chemical compounds, when applied in a somewhat different way, are the general considerations that govern the magnitude of the inter-atomic distance in both elements and compounds. Indeed, all aggregates of electronegative elements *are* molecular in their composition, rather than atomic, as the molecular requirement that the negative electric displacement of an atom of such an element must be counterbalanced by an equivalent positive displacement in order to arrive at a stable equilibrium in space applies with equal force to a combination with a like atom. As we saw in our examination of the structural situation, electropositive elements are not subject to this restriction, but in many cases the molecular (balanced orientation) type of structure takes precedence over the electropositive structure by reason of collateral factors that affect the relative probability. Because of this fact that the distances follow the structural pattern, the various ways of orienting the atomic rotations that were discussed in Chapter 18, Vol. I, with a few modifications due to the special conditions that exist in the elemental aggregates, determine the manner in which the atoms of an element are able to combine with each other, and the effective values of the specific rotations in these combinations.

In the electropositive elements the specific rotations are based, in the first instance, on the rotational displacements as listed in Chapter 10, Vol. I. Where the inter-atomic orientation is the normal positive arrangement, the displacements as listed are translated directly into specific rotations by addition of the initial unit and reduction of the incremental values where the rotation extends to vibration two. Except for the elements of group 2A, which, as already noted, are subject to some special considerations because of their low magnetic displacements, the elements of Division I all follow the regular electropositive pattern of specific rotations. The only irregularities are in the electric rotations

of the second and third elements of each group, where the point of transition to vibration two varies between groups. The inter-atomic distances in this division are listed in Table 2.

TABLE 2: *DISTANCES-DIVISION I*

Group	Atomic Number	Element	Specific Rotation		Distance	
			Magnetic	Electric	Calc.	Obs.
2B	11	Sodium	3-2½	2	3.70	3.71
			3-3			
	12	Magnesium	3-2½	2½	3.17	3.21
3A	13	Aluminum	3-2½	3	2.83	2.86
	19	Potassium	4-3	2	4.49	4.50
	20	Calcium	4-3	2½	4.00	3.98
	21	Scandium	4-3	4	3.18	3.20
	22	Titanium	4-3	5	2.95	2.92
3B	37	Rubidium	4-4	2	4.85	4.87
	38	Strontium	4-4	2½	4.32	4.28
	39	Yttrium	4-4	3½	3.64	3.63
	40	Zirconium	4-4	5	3.18	3.23
4A	55	Cesium	4½-4½	2	5.23	5.24
	56	Barium	5-4½	3	4.36	4.34
	57	Lanthanum	4½-4½	4	3.70	3.74
	58	Cerium	5-4½	5	3.61	3.63
4B	89	Actinium	4½-5	4	3.79	3.76*
	90	Thorium	4½-5	5	3.52	3.56

The regular electropositive pattern is also applicable in Division II, and a number of the Division II elements of Group 3A crystallize on this basis, with inter-atomic distances determined in the same manner as in Division I. As noted in Volume I, however, the Division II elements generally favor the magnetic type of orientation in chemical compounds because the normal positive orientation becomes less probable as the displacement increases. The same probability considerations operate against the positive orientation in the elements of this division, but instead of employing the magnetic orientation as the alternate, these elements utilize a type of orientation that is available only where all rotations of each participant in a combination are identical with those of the other. This arrangement reverses the effective directions of the rotations of alternate atoms. The resulting

relative rotation is a combination of x and $8-x$ (or $4-x$), as in the neutral orientation, and the effective specific rotations are 10 for vibration one and 5 for vibration two. A combination value 5-10 is also common.

This reverse type of structure makes its appearance in body-centered cubic crystal forms of chromium and iron which coexist with the regular positive hexagonal or face-centered cubic structures. Vanadium and niobium, the first Division II elements of their respective groups, combine the positive and reverse orientations. Beyond niobium the positive orientation does not appear in the common Division II forms of the elements, the structures to which the present discussion is limited, and all elements take the reverse orientation, except europium and ytterbium, which combine it with a unit-specific rotation; that is, no electric-rotational displacement at all, as in the inert gas elements.

On the basis of the considerations discussed in Chapter 1, the average effective specific rotation for such rotational combinations has been taken as the geometric mean of the two components. Where the orientations are the same, and the only difference is in the magnitude, as in the 5-10 combination, and in the combinations of magnetic rotations that we will encounter later, the equilibrium is reached in the normal manner. If two different electric rotations are involved, the two-atom pairs cannot attain spatial equilibrium individually, but they establish a group equilibrium similar to that which is achieved where n atoms of valence one each combine with one atom of valence n .

The Division II distances are shown in Table 3. Because of the greater probability of the electropositive types of combinations, the characteristics of Division II carry over into the first elements of Division III, and these elements, nickel, palladium, and lutetium, are included in the table. Some similar modifications of the normal division boundaries have already been noted in connection with other subjects.

The net total rotation of the material atom is a motion with positive displacement—that is, a speed less than unity—and as such it normally results in a change of position in space. Inside unit space, however, all motion is in time. The orientation of the atom for the purpose of the space-time equilibrium therefore exists in the three dimensions of time. As we saw in our examination of the inter-regional situation in Chapter 12, Volume I, each of these dimensions contacts the space of the region outside unit distance individually. To the extent that

TABLE 3: *DISTANCES-DIVISION II*

Group	Atomic Number	Element	Specific Rotation		Distance	
			Magnetic	Electric	Calc.	Obs.
3A	23	Vanadium	4-3	6-10	2.62	2.62
	24	Chromium	4-3	7	2.68	2.72
			4-3	10	2.46	2.49
	25	Manganese	4-3	8	2.59	2.58
	26	Iron	4-3	8½	2.56	2.57
			4-3	10	2.46	2.48
	27	Cobalt	4-3	9	2.52	2.51
	28	Nickel	4-3	9½	2.49	2.49
	41	Niobium	4-4	6-10	2.83	2.85
	42	Molybdenum	4-4½	10	2.72	2.72
3B	43	Technetium	4-4½	10	2.73	2.73*
	44	Ruthenium	4-4½	10	2.73	2.70
	45	Rhodium	4-4	10	2.66	2.69
			4-4½	10	2.73	2.76
	46	Palladium	4-4½	10	2.73	2.74
	59	Praseodymium	5-4½	5	3.61	3.64
	60	Neodymium	5-4½	5	3.61	3.65
	62	Samarium	5-4½	5	3.61	3.62*
	63	Europium	4½-5	1-5	3.96	3.96
	64	Gadolinium	5-4½	5	3.61	3.62
4A	65	Terbium	5-4½	5	3.61	3.59
	66	Dysprosium	5-4½	5	3.61	3.58
	67	Holmium	4½-5	5	3.52	3.56
	68	Erbium	4½-5	5	3.52	3.53
	69	Thulium	4½-5	5	3.52	3.52
	70	Ytterbium	4½-4½	1-5	3.86	3.87
	71	Lutetium	4½-5	5	3.52	3.50*
	91	Protactinium	4½-5	5-10	3.22	3.24*
	92	Uranium	4½-4½	10	2.87	2.85
	93	Neptunium	4½-4½	5	3.43	3.46*
4B	94	Plutonium	4½-4½	5-10	3.14	3.15*
	95	Americum	4½-4½	5	3.43	3.46*
	96	Curium	4½-4½	5-10	3.14	3.10*
	97	Berkelium	4½-4½	5	3.43	3.40*

the motion in a dimension of time acts along the line of this contact it is a motion in *equivalent space*. Otherwise it has no spatial effect beyond the unit boundary. Because of the independence of the three dimensions of motion in time the relative orientation of the electric rotation of any combination of atoms may be the same in all spatial dimensions, or there may be two or three different orientations.

In most of the elements that have been discussed thus far the orientation is the same in all spatial dimensions, and in the exceptions the alternate rotations are symmetrically distributed in the solid structure. The force system of an aggregate of such elements is isotropic. It follows that any aggregate of atoms of these elements has a structure in which the constituents are arranged in one of the geometrical patterns possible for equal forces: an isometric crystal. All of the electropositive elements (Divisions I and II) crystallize in isometric forms, and, except for a few which apparently have quite complex structures, each of the crystal forms of these elements belongs to one or another of three types: the face-centered cube, the body-centered cube, or the hexagonal close-packed structure.

We now turn to the other major subdivision of the elements, the electronegative class, those whose normal electric displacement is negative. Here the force system is not necessarily isotropic, since the most probable arrangement in one or two dimensions may be the *negative orientation*, a direct combination of two negative electric displacements, similar to the all-positive combinations. It is not possible to have negative orientation in all three dimensions, and wherever it does exist in one or two dimensions the rotational forces of the atoms are necessarily anisotropic. The controlling factor is the requirement that the net total rotational displacement of a material atom as a whole must be positive. Negative orientation in all three dimensions is obviously incompatible with this requirement, but if the negative displacement is restricted to one dimension the aggregate has fixed atomic positions in two dimensions, with a fixed average position in the third because of the positive displacement of the atom as a whole. This results in a crystal structure that is essentially equivalent to one with fixed positions in all dimensions. Such crystals are not usually isometric, as the inter-atomic distance in the odd dimension is generally different from that of the other

two. Where the distances in all dimensions do happen to coincide, we will find on further investigation that the space symmetry is not an indication of force symmetry.

If the negative displacement is very small, as in the lower division IV elements, it is possible to have negative orientation in two dimensions if the positive displacement in the third dimension exceeds the sum of these two negative components, so that the net result is still positive. Here the relative positions of the atoms are fixed in one dimension only, but the average positions in the other two dimensions are constant by reason of the net positive displacement of the atoms. An aggregate of such atoms retains most of the external characteristics of a crystal, but when the internal structure is examined the atoms appear to be distributed at random, rather than in the orderly arrangement of the crystal. In reality there is just as much order as in the crystalline structure, but part of the order is in time rather than in space. This form of matter can be identified as the *glassy*, or *vitreous*, form, to distinguish it from the *crystalline* form.

The term "state" is frequently used in this connection instead of "form," but the physical state of matter has an altogether different meaning based on other criteria, and it seems advisable to confine the use of this term to the one application. Both glasses and crystals are in the *solid* state.

In beginning a consideration of the structures of the individual electronegative elements, we will start with Division III. The general situation in this division is similar to that in Division II, but the negativity of the normal electric displacement introduces a new factor into the determination of the orientation pattern, as the most probable orientation of an electronegative element may not be capable of existing in all three dimensions. As stated earlier, where two or more different orientations are possible in a given set of circumstances the relative probability is the deciding factor. Low displacements are more probable than high displacements. Simple orientations are more probable than combinations. Positive electric orientation is more probable than negative. In Division I all of these factors operate in the same direction. The positive orientation is simple, and it also has the lowest displacement value. All structures in this

division are therefore formed on the basis of the positive orientation. In Division II the margin of probability is narrow. Here the positive displacement x is greater than the inverse displacement $\delta-x$, and this operates against the greater inherent probability of a simple positive structure. As a result, both the positive and reverse types of structure are found in this division, together with a combination of the two.

In Division III the negative orientation has a status somewhat similar to that of the positive orientation in Division II. As a simple orientation, it has a relatively high probability. But it is limited to one dimension. The regular Division III structures of Groups 3A and 3B are therefore anisotropic, with the reverse orientation in the other two dimensions. A combination of these two types of orientation is also possible, and in copper and silver, the first Division III elements of their respective groups, the crystals formed on the basis of this combination orientation have cubic symmetry. As in Division II, the elements of Division III in Groups 4A and 4B crystallize entirely on the basis of the reverse orientation. Table 4 lists what may be considered as the regular inter-atomic distances of the elements of Division III.

Although the probability of the negative orientation is greater in Division IV than in Division II, because of the smaller displacement values, this type of structure seldom appears in the crystals of the lower division. The reason is that where this orientation exists in the elements of the lower displacements, it exists in two dimensions, and this produces a glassy or vitreous aggregate rather than a crystal. The reverse orientation is not subject to any restrictive factor of this nature, but it is less probable at the lower displacements, and except in Group 4A, where it continues to predominate, this orientation appears less frequently as the displacement decreases. Where it does exist it is increasingly likely to combine with some other type of orientation. As a result of these limitations that are applicable to the inherently more probable types of orientation, many of the Division IV structures are formed on the basis of the *secondary positive* orientation, a combination of two $\delta-x$ displacements.

The secondary positive orientation is not possible in the electro-positive divisions, as $\delta-x$ is negative in these divisions, and like

the negative orientation itself, an 8-x negative combination would be confined to a subordinate role in one or two dimensions of an asymmetric structure. Such a crystal structure cannot compete with the high probability of the symmetrical electropositive crystals, and therefore does not exist. In the electronegative divisions, however, the 8-x displacement is positive, and there are no limitations on it, aside from those arising from the high displacement values.

TABLE 4: *DISTANCES-DIVISION III*

Group	Atomic Number	Element	Specific Rotation		Distance	
			Magnetic	Electric	Calc.	Obs.
3A	29	Copper	4-3	8-10	2.53	2.55
	30	Zinc	4-4	7	2.90	2.91
			4-4	10	2.66	2.66
			4-3	6	2.79	2.80
	31	Gallium	4-3	10	2.46	2.44
3B	47	Silver	4-5	8-10	2.87	2.88
	48	Cadmium	5-4	7	3.20	3.26*
			5-4	10	2.94	2.97
	49	Indium	5-4	6	3.33	3.37
			5-4	6-10	3.21	3.24
4A	72	Hafnium	4-4½	5	3.26	3.32
	73	Tantalum	4½-4½	10	2.87	2.86
	74	Tungsten	4-4½	10	2.73	2.74
	75	Rhenium	4-4½	10	2.73	2.77*
	76	Osmium	4-4½	10	2.73	2.73
	77	Iridium	4-4½	10	2.73	2.71
	78	Platinum	4-4½	10	2.73	2.77
	79	Gold	4½-4½	10	2.87	2.88
	80	Mercury	4-4½	5-10	2.98	3.00
			4½-4½	5	3.43	3.47
	81	Thallium	4½-4½	5	3.43	3.45

The effective displacement of this secondary positive orientation is even greater than might be expected from the magnitude of the quantity 8-x, as the change of zero points for the two oppositely directed motions is also oppositely directed, and the new zero points are 16 displacement units apart. The resultant relative displacement is 16-2x, and the corresponding specific rotation is 18-2x. In

Division IV the numerical values of the latter expression range from 10 to 16, and because of the low probability of such high rotations, the secondary positive orientation is limited to one or one and one-half dimensions in spite of its positive character. In Division III the 8-x displacements are lower, but in this case they are too low. A two-unit separation of the zero points (16 displacement units) cannot be maintained unless the effective displacement is at least 8 (one full three-dimensional unit). The secondary positive orientation is therefore confined to Division IV.

A special type of structure is possible only for those electronegative elements which have a rotational displacement of four units in the electric dimension. These elements are on the borderline between Divisions III and IV, where the secondary positive and reverse orientations are about equally probable. Under similar conditions other elements crystallize in hexagonal or tetragonal structures, utilizing the different orientations in the different dimensions. For these displacement 4 elements, however, the two orientations produce the same specific rotation: 10. The inter-atomic distance in these crystals is therefore the same in all dimensions, and the crystals are isometric, even though the rotational forces in the different dimensions are not of the same character. The molecular arrangement in this crystal pattern, the *diamond* structure, shows the true nature of the rotational forces. Outwardly this crystal cannot be distinguished from the isotropic cubic crystals, but the analogous body-centered cubic structure has an atom at each corner of the cube as well as one in the center, whereas the diamond structure leaves alternate corners open to accommodate the abnormal projection of forces in the secondary positive dimension.

In those of the lower elements of Division IV that are beyond the range of the inverse type of orientation, there is no available alternative for combination with the secondary positive orientation. The crystals of these elements therefore have no effective electric rotation in the remaining dimensions, and the relative specific rotation in these dimensions is unity, as in all dimensions of the inert gas elements. The most common distances in the aggregates of the Division IV elements are shown in Table 5.

TABLE 5: *DISTANCES-DIVISION IV*

Group	Atomic Number	Element	Specific Rotation		Distance	
			Magnetic	Electric	Calc.	Obs.
2B	14	Silicon	3-3	5-10	2.31	2.35
			3-3	10	2.19	2.20
	15	Phosphorus	3-4	1	3.46	3.48*
			3-3	10	2.11	2.07
			3-3	1	3.21	3.27*
			3-3	16	1.92	1.82
	16	Sulfur	3-3	1-16	2.48	2.52
			4-3	10	2.46	2.43
3A	32	Germanium	4-3	12	2.37	2.44*
			4-3	10	2.46	2.51
	33	Arsenic	4-3	14	2.32	2.32
			3-4	1	3.46	3.46
	34	Selenium	4-3	16	2.25	2.27
			3-4	1	3.46	3.30
	35	Bromine	4-3	10	2.80	2.80
			5-4	5-10	3.22	3.17
3B	50	Tin	5-4	10	2.94	3.02
			5-4	12	2.83	2.87
	51	Antimony	5-4	4-10	3.34	3.36*
			5-4½	14	2.82	2.86
	52	Tellurium	5-4½	1-10	3.71	3.74
			5-4	16	2.68	2.70
	53	Iodine	5-4	1-16	3.54	3.54
			5-4	1	4.46	4.41*
4A	82	Lead	4½-4½	5	3.43	3.49
			4½-4½	5	3.43	3.47*
	83	Bismuth	4½-4½	5-10	3.14	3.10
			4½-4½	5	2.43	3.40*

Up to this point, no consideration has been given to the elements of atomic number below 10, as the rotational forces of these elements are subject to certain special influences which make it desirable to discuss them separately. One cause of deviation from the normal behavior is the small size of the rotational groups. In the larger groups the four divisions are distinct, and, except for some

overlapping, each has its own characteristic force combinations, as we have seen in the preceding paragraphs. In an 8-element group, however, the second series of four elements, which would normally constitute Division II, is actually in the Division IV position. As a result, these four elements have, to a certain extent, the properties of both divisions. Similarly, the Division I elements of these groups may, in some cases, act as if they were members of Division III.

A second influence that affects the forces and the crystal structures of the lower group elements is the inactivity of the rotational forces in certain dimensions that was mentioned earlier. A specific rotation of two units produces no effect in the positive direction. The reason for this is revealed by equation 1-1. By applying this equation we find that the effective rotational force ($\ln t$) for $t=2$ is 0.693, which is less than the opposing space-time force 1.00. The net effective force of specific rotation 2 is therefore below the minimum value for action in the positive direction. In order to produce an active force the specific rotation must be high enough to make $\ln t$ greater than unity. This is accomplished at rotation 3.

The specific magnetic rotation of the 1B group, which includes only the two elements hydrogen and helium, and the 2A group of eight elements beginning with lithium, combines the values 3 and 2. Where the value 2 applies to the subordinate rotation (3-2), one dimension is inactive; where it applies to the principal rotation (2-3), two dimensions are inactive. This reduces the force exerted by each atom to $\frac{2}{3}$ of the normal amount in the case of one inactive dimension, and to $\frac{1}{3}$ for two inactive dimensions. The inter-atomic distance is proportional to the square root of the product of the two forces involved. Thus the reduction in distance is also $\frac{1}{3}$ per inactive dimension.

Since the electric rotation is not a basic motion, but a reverse rotation of the magnetic rotational system, the limitations to which the basic rotation is subject are not applicable. The electric rotation merely modifies the magnetic rotation, and the low value of the force integral for specific rotation 2 makes itself apparent by an inter-atomic distance which is greater than that which would prevail if there were no electric displacement at all (unit specific rotation).

Theoretical values of the inter-atomic distances of the lower group

elements are compared with measured values in Table 6. The figures in parentheses in column 4 of this table indicate the effective number of dimensions. Thus the notation 3(1) shown for hydrogen means that this element has a specific magnetic rotation of 3, effective in only one dimension.

TABLE 6: *DISTANCES-LOWER GROUP ELEMENTS*

Group	Atomic Number	Element	Specific Rotation		Distance	
			Magnetic	Electric	Calc.	Obs.
1B	1	Hydrogen	3(1)	10	0.70	0.74*
	2	Helium	3(1)	1	1.07	1.09*
2A	3	Lithium	2½-2½	2	3.05	3.03
	4	Beryllium	3(2)	2½	2.28	2.28
	5	Boron	3(2)	5	1.68	1.74*
	5	Boron	3-3	10	2.11	2.03*
	6	C (diamond)	3(2)	5-10	1.54	1.54
	6	C (graphite)	3(2)	10	1.41	1.42
	6	C	3-3	1	3.21	3.40
	7	Nitrogen	3(1½)	10	1.06	1.06
	7	Nitrogen	3-3	1	3.21	3.44*
	8	Oxygen	3(1½)	10	1.06	1.15*
	8	Oxygen	3-3	1	3.21	3.2 *
	9	Fluorine	3(2)	10	1.41	1.44*

Except where the crystals are isometric, there is still much uncertainty in the distance measurements on these lower group elements, and many other values have been reported in addition to those included in the table. This situation will be discussed at length in Chapter 3, where we will have the benefit of measurements of the distances between like atoms that are constituents of chemical compounds.

As indicated in the introductory paragraphs of this chapter, we are not yet in a position where we can determine specifically just what the inter-atomic distance will be for any given element under a given set of conditions. The theoretical considerations that have been discussed actually do lead to specific values in many cases, but in other instances there is an uncertainty as to which of two or more theoretically possible rotational arrangements corresponds to the observed crystal structure. Continuing progress is being made

in both the experimental and the theoretical fields, and it can be expected that these uncertainties will gradually diminish toward the irreducible minimum that was mentioned earlier. In the course of this process there will necessarily be some changes in the identifications of the observed inter-atomic distances with the theoretically possible structures. A comparison of Tables 1 to 6 with the corresponding tabulations of the first edition should therefore be of interest as an indication of the nature and magnitude of the changes that have taken place in our view of this inter-atomic distance situation in the last twenty years, and by extension, an indication of the amount of change that can be expected in the future.

Such a comparison shows that the modifications of the original conclusions that now appear to be required, in the light of the additional information that has been made available, are confined almost entirely to those which have resulted from a better theoretical understanding of the behavior of the specific magnetic rotation above an effective value of 4. Few changes are required in either the magnetic or electric values in those rotational combinations where the specific magnetic rotation is 4-4 or less.

One of the puzzling features of the rotational situation as it appeared at the time of the original publication was the apparent retrograde progression of the specific magnetic rotation in Groups 4A and 4B. It was recognized at that time that both the $4\frac{1}{2}$ and 5 values of the specific rotation correspond to the same displacement, 4, the difference being that in the case of the $4\frac{1}{2}$ value the rotation extends to two units of vibration, and the last increment of specific rotation in this case is only half size. The next half unit increment, if such an increment were possible, would bring the $4\frac{1}{2}$ rotation back to the 5 value. It would therefore appear that the sequence of specific rotations beyond $4\frac{1}{2}$ -4 should be $4\frac{1}{2}$ - $4\frac{1}{2}$, 5- $4\frac{1}{2}$, 5-5, and so on. But the tendency is in the opposite direction. Instead of moving toward higher values as the atomic number increases, there is actually a *decreasing* trend. This was already evident at the time of publication of the first edition, as the low inter-atomic distances of the series of elements from tungsten to platinum could not be accounted for unless the specific magnetic rotation dropped back to 4- $4\frac{1}{2}$ from the higher levels of the preceding elements of the 4A group. This decreasing trend has become even

more prominent as distances have become available for additional elements of Group 4B, as some of these values indicate specific magnetic rotations of 4-4, or possibly even 4-3½.

As it happens, the continuation of the trend toward lower values in the more recent data has had the effect of clarifying the situation. It is now evident that the 5-5 specific rotation is not reached within the accessible portion of Groups 4A and 4B. (Considerations that will be discussed later show that the specific rotation of 5-5 would be unstable.) The lower values in the 4A and 4B groups do not result from a decrease in the magnetic displacement, but from a shift of the existing displacement units from vibration one to vibration two, a process which reduces the specific rotation of the units by one half. On a vibration one basis, rotational displacements 4-3 correspond to specific rotations 5-4. Conversion of successive units of displacement to vibration two, without change in the number of displacement units, results in a series of specific rotations, 5-4, 4½-4, 4-4½, 4-4, and so on. A similar series with one additional displacement unit goes through the values 5-4½, 4½-5, 4½-4½, 4½-4, and then follows the same route as the series with the lower displacement.

The modifications that have been made in the theoretical rotational values applicable to the elements of these two highest rotational groups since the publication of the first edition are the result of a review of the situation in the light of this new understanding of the trend of the specific rotation. The general pattern in group 4A is now seen to be that of the series from 5-4½ to 4-4½, with a return to 4½-4½ in the lower electronegative elements. So far as can be determined at this time, Group 4B follows the same pattern one step farther advanced; that is, it begins with 4½-5 rather than 5-4½.

The difference in the inter-atomic distance corresponding to one of the steps in this conversion process is relatively small, and in view of the substantial variation in the experimental values it has not appeared advisable to take into account the possibility of combinations such as 4½-5 specific rotation of one atom of a pair and 4½-4½ in the other. It seems clear that such combinations do exist in some of the lower group elements, sodium, for example, and they probably play some part in the higher groups. Most of

the reported distances for holmium and erbium, for instance, agree more closely with a combination of $5-4\frac{1}{2}$ and $4\frac{1}{2}-5$ than with either individually. However, all of these values are theoretically possible, and the only question at issue in this and many other similar cases is *which* theoretical value corresponds to the observed distance. Definitive answers to identification questions of this kind will have to wait until the theoretical probabilities are specifically evaluated, or the experimental uncertainties are resolved.

Many questions concerning alternate crystal structures will also have to wait for more information from theory or experiment, particularly where crystal forms that exist only at high temperatures or pressures are involved. There is, however, a large body of information already available in this area, and it can be tied into the theoretical picture as soon as someone has the time and the inclination to undertake the task.

CHAPTER 3

Distances in Compounds

Thus far in the discussion of the inter-atomic distances we have been dealing with aggregates composed of like atoms. The same general principles apply to aggregates of unlike atoms, but the existence of differences between the components of such systems introduces some new factors that we will now want to examine.

The matters to be considered in this chapter have no relevance to direct combinations of electropositive elements (aggregates of which are mixtures or alloys, rather than chemical compounds). As noted in Chapter 18, Vol. I, the proportions in which such elements can combine may be determined, or limited, by geometrical considerations, but aside from such effects, unlike atoms of this kind can combine on the same basis as like atoms. Here the forces are identical in character and concurrent, the type of combination that we have called the positive orientation. The resultant specific electric rotation, according to the principles previously set forth, is $(t_1 \times t_2)^{1/2}$, the geometric mean of the two constituents. If the two elements have different magnetic rotations, the resultant is also the geometric mean of the individual rotations, as the magnetic rotations always have positive displacements, and these combine in the same manner as the positive electric displacements. The effective electric and magnetic specific rotations thus derived can then be entered in the applicable force and distance equations from Chapter 1.

Combinations of unlike positive atoms may also take place on the basis of the reverse orientation, the alternate type of structure that is available to the elemental aggregates. Where the electric rotations of the components differ, the resultant specific rotation of the two-atom combination will not be the required neutral 5 or 10, but a second pair of atoms inversely oriented to the first results in a four-atom group that has the necessary rotational balance.

As brought out in Volume I, the simplest type of combination in chemical compounds is based on the normal orientation, in which Division I electropositive elements are joined with Division IV electronegative elements on the basis of numerically equal displacements. The resultant effective specific magnetic rotation can be calculated in the same manner as in the all-positive structures, but, as we saw in our consideration of the inter-atomic distances of the elements, where an equilibrium is established between positive and negative electric rotations, the resultant is the sum of the two individual values, rather than the mean.

When this arrangement unites one electropositive atom with each electronegative atom the resulting structure is usually a simple cube with the atoms of each element occupying alternate comers of the cube. This is called the *sodium chloride* structure, after the most familiar member of the family of compounds crystallizing in this form. Table 7 gives the inter-atomic distances of a number of common NaCl type crystals.

TABLE 7: *DISTANCES-NA CL TYPE COMPOUNDS*

Compound	Specific Rotation		Elec.	Distance	
		Magnetic		Calc.	Obs.
LiH	3(2)	3(2)	3	2.04	2.04
LiF	3(2)	3(2)	3	2.04	2.01
LiCl	3(2)	3½-3½	4	2.57	2.57
LiBr	3(2)	4-4	4	2.77	2.75
Li	3(2)	5-4	4	2.96	3.00
NaF	3-2½	3(2)	4	2.26	2.31
NaCl	3-2½	3½-3½	4	2.77	2.81
NaBr	3-2½	4-4	4	2.94	2.98
NaI	3-3	5-4	4	3.21	3.23
MgO	3-3	3(2)	5½	2.15	2.10
MgS	3-3	3½-3½	5½	2.60	2.59
MgSe	3-3	4-4	5½	2.76	2.72
KF	4-3	3(2)	4	2.63	2.67
KCl	4-3	3½-3½	4	3.11	3.14
KBr	4-3	4-4	4	3.30	3.29
KI	4-3	5-4	4	3.47	3.52

Compound	Specific Rotation			Distance	
	Magnetic		Elec.	Calc.	Obs.
CaO	4-3	3(2)	5½	2.38	2.40
CaS	4-3	3½-3½	5½	2.81	2.84
CaSe	4-3	4-4	5½	2.98	2.95
CaTe	4-3	5-4	5½	3.13	3.17
ScN	4-3	3(2)	7	2.22	2.22
TiC	4-3	3(2)	8½	2.12	2.16
RbF	4-4	3(2)	4	2.77	2.82
RbCl	4-4	3½-3½	4	3.24	3.27
RbBr	4-4	4-4	4	3.43	3.43
RbI	4-4	5-4	4	3.61	3.66
SrO	4-4	3(2)	5½	2.51	2.57
SrS	4-4	3½-3½	5½	2.92	2.93
SrSe	4-4	4-4	5½	3.10	3.11
SrTe	4-4	5-4	5½	3.26	3.24
CsF	5-4	3(2)	4	2.96	3.00
CsCl	5-4	4-3	4	3.47	3.51
BaO	5-4½	3(2)	5½	2.72	2.76
BaS	5-4½	4-3	5½	3.17	3.17
BaSe	5-4½	4-4	5½	3.30	3.31
BaTe	5-4½	5-4	5½	3.47	3.49
LaN	5-4	3(2)	6	2.61	2.63
LaP	5-4	4-3	6½	2.99	3.01
LaAs	5-4	4-4	7	3.04	3.06
LaSb	5-4	5-4	7	3.20	3.24
LaBi	5-4	5-4½	7	3.24	3.28

From this tabulation it can be seen that the special rotational characteristics which certain of the elements possess in the elemental aggregates carry over into their compounds. The second element in each group shows the same preference for rotation on the basis of vibration two that we encountered in examining the structures of the elements. Here, again, this preference extends to some of the following elements, and in such series of compounds as CaO, ScN, TiC, one component keeps the vibration two status throughout the series, and the resulting effective rotations are 5½, 7, 8½, rather than 6, 8, 10. The elements of the lower groups have inactive force dimensions in the compounds just as in the elemental structures previously examined.

If the active dimensions are not the same in both components, the full rotational force of the more active component is effective in its excess dimensions, the effective rotation in an inactive dimension being unity. For example, the value of $\ln t$ for magnetic rotation 3 is 1.099 in three dimensions, or 0.7324 in two dimensions. If this two-dimensional rotation is combined with a three-dimensional magnetic rotation x , the resultant value of $\ln t$ is $(0.7324 \ x)^{1/2}$, the geometric mean of the individual values, in two dimensions, and x in the third. The average value for all three dimensions is $(0.7324 \ x^2)^{1/3}$.

This dimensional inactivity in the lower groups plays only a minor role in the structures of the elements, as can be seen from the fact that it did not need any attention until almost the end of Chapter 2. In the compounds, however, it is very significant, because the compounds that contain lower group elements (below atomic number 10) constitute the great bulk of all chemical compounds.

Except for certain types of crystals that are essentially interchangeable, the structures of the elements are determined almost entirely by the nature of the orientations. In compounds there is another active factor: the relative proportions of the components. Where two atoms of one kind form a compound with one atom of another on the basis of the normal orientation, the unequal proportions make the NaCl arrangement impossible, and instead the crystal has the *calcium fluoride* structure, which is also cubic but has a different atomic arrangement. Inter-atomic distances for a number of common CaF_2 type crystals are listed in Table 8.

TABLE 8: *DISTANCES- CaF_2 TYPE COMPOUNDS*

Compound	Specific Rotation			Distance	
	Magnetic		Electric	Calc.	Obs.
Na_2O	3-2½	3(2)	3½	2.39	2.40
Na_2S	3-2½	4-3	4	2.83	2.83
Na_2Se	3-2½	4-4	4	2.94	2.95
Na_2Te	3-2½	5-4½	4	3.13	3.17
Mg_2Si	3-3	4-3	5	2.73	2.77
Mg_2Ge	3-3	4-4	5½	2.76	2.76
Mg_2Sn	3-3	5-4	5½	2.90	2.93
Mg_2Pb	3-3	5-4½	5½	2.94	2.96
K_2O	4-3	3(2)	3½	2.79	2.79

Compound	Specific Rotation		Distance		
	Magnetic	Electric	Calc.	Obs.	
K ₂ S	4-3	4-3	4	3.17	3.20
K ₂ Se	4-3	4-4	4	3.30	3.33
K ₂ Te	4-3	5-4½	4	3.51	3.53
CaF ₂	4-3	3(2)	5½	2.38	2.36
Rb ₂ O	4-4	3(2)	3½	2.94	2.92
Rb ₂ S	4-4	4-3	4	3.30	3.31
SrF ₂	4-4	3(2)	5½	2.50	2.50
SrCl ₂	4-4	4-3	5½	2.98	3.03
BaF ₂	5-4	3(2)	5½	2.68	2.68
BaCl ₂	5-4½	4-3	5½	3.17	3.18*

The compounds of lithium with valence one negative elements follow the regular pattern, and were included in Table 7, but the compounds with valence two elements are irregular, and they have therefore been omitted from Table 8. As we will see in Chapter 6, the irregularity is due to the fact that the two lithium atoms in a molecule of the CaF₂ type act as a radical rather than as independent constituents of the molecule.

These two normal orientation tables, 7 and 8, provide an impressive confirmation of the validity of the theoretical findings. One of the problems in dealing with the inter-atomic distances of the elements is that because of the relatively small total number of elements, the number to which any particular magnetic rotational combination is applicable is quite small, and consequently it is rather difficult to establish a *prima facie* case for the authenticity of the rotational values. But this is not true of the normal type compounds, as they are more numerous and less variable. There are two elements in these tables, sulfur and chlorine, that have different magnetic rotations under different conditions. These elements have 4-3 rotation in the CaF₂ type crystals, and in the NaCl type combinations with elements of group 4A. In the other compounds of the NaCl type they take the 3½-3½ rotations. There are also two more elements, each of which, according to the information now available, deviates from its normal rotations in one of the listed compounds. Otherwise, all of the elements entering into the 60 compounds in the two tables have the same specific magnetic rotations in every compound in which they participate.

Furthermore, when the inherent differences between the elemental and compound aggregates are taken into account, there is also agreement between these rotations in the compounds and the specific rotations of the same elements in the elemental aggregates. The most common difference of this kind is a result of the fact that the Division IV element in a compound has a purely negative role. For this reason it takes the magnetic rotation of the next higher group. In the elemental aggregates half of the atoms are reoriented to act in a positive capacity. Consequently, they tend to retain the normal rotation of the group to which they actually belong. For example, the Division IV elements of Group 3A, germanium, arsenic, selenium, and bromine, have the normal specific rotation of their group, 4-3, in the crystals of the elements, but in the compounds they take the 4-4 specific rotation of Group 3B, acting as negative members of that group.

Another difference between the two classes of structures is that those elements of the higher groups that have the option of extending their rotation to a second vibrational unit are less likely to do so if they are combining with an element which is rotating entirely on the basis of vibration one. Aside from these deviations due to known causes, the values of the specific magnetic rotation determined for the elements in Chapter 2 are also generally applicable to the compounds. This equivalence does not apply to the specific electric rotations, as they are determined by the way in which the rotations of the constituents of each aggregate are oriented relative to each other, a relation that is different in the two classes of structures.

This applicability of the *same* equations and, in general, the *same* numerical values, to the calculation of distances in both elements and compounds contrasts sharply with the conventional theory that regards the inter-atomic distance as being determined by the “sizes” of the atoms. The sodium atom, or “ion,” in the NaCl crystal, for example, is asserted to have a radius only about 60 percent as large as the radius of the atom in the elemental aggregate. If this atom takes part in a compound which cannot be included in the “ionic” class, current theory gives it still a different “size”: what is called a “covalent” radius. The need for assuming any extraordinary changeability in the size of what, so far as we can tell, continues to be the same object, is now eliminated by the finding that the variations in the inter-atomic

distance have nothing to do with the sizes of the atoms, but merely indicate differences in the location of the equilibrium between the inward and the outward forces to which the atoms are subject.

Another type of orientation that forms a relatively simple binary compound is the rotational combination that we found in the diamond structure. As in the elements, this is an equilibrium between an atom of a Division IV element and one of Division III, the requirement being that $t_1 + t_2 = 8$. Obviously, the only elements that can meet this requirement by themselves are those whose negative rotational displacement (valence) is 4, but any Division IV element can establish an equilibrium of this kind with an appropriate Division III element.

Closely associated with this cubic diamond-like *zinc sulfide* class of crystals is a hexagonal structure based on the same orientation, and containing the same equal proportions of the two constituents. Since these controlling factors are identical in the two forms, the crystals of the hexagonal *zinc oxide* class have the same inter-atomic distances as the corresponding zinc sulfide structures. In such instances, where the inter-atomic forces are the same, there is little or no probability advantage of one type of crystal over the other, and either may be formed under appropriate conditions. Table 9 lists the inter-atomic distances for some common crystals of these two classes.

TABLE 9: *DISTANCES-DIAMOND TYPE COMPOUNDS*

Compound	Specific Rotation			Distance	
	Magnetic		Electric	Calc.	Obs.
ZnS (Cubic) Class					
AlP	3-4	3½-3½	10	2.32	2.35
AlAs	3-4	4-4	10	2.62	2.43
AlSb	3-4	5-4½	10	2.62	2.66
SiC	3-4	3(2)	10	1.94	1.93*
CuCl	3-4	3½-3½	10	2.32	2.35
CuBr	3-4	4-4	10	2.46	2.46
CuI	3-4	5-4	10	2.59	2.62
ZnS	3-4	3½-3½	10	2.32	2.36
ZnSe	3-4	4-4	10	2.46	2.45
ZnTe	3-4	5-4½	10	2.62	2.63*
GaP	3-4	3½-3½	10	2.32	2.36
GaAs	3-4	4-4	10	2.46	2.43

Compound	Specific Rotation			Distance	
	Magnetic		Electric	Calc.	Obs.
ZnS (Cubic) Class (cont'd)					
GaAs	3-4	4-4	10	2.46	2.43
GaSb	3-4	5-4½	10	2.62	2.65
AgI	4-4	5-4	10	2.80	2.81
CdS	4-4	3½-3½	10	2.51	2.52
CdTe	4-4	5-4	10	2.80	2.78
InP	4-4	3½-3½	10	2.51	2.54
InAs	4-4	4-4	10	2.66	2.62
InSb	4-4	5-4	10	2.80	2.80
ZnO (Hexagonal) Class					
AlN	3-4	3(2)	10	1.94	1.90
ZnO	3-4	3(2)	10	1.94	1.95
ZnS	3-4	3½-3½	10	2.32	2.33
GaN	3-4	3(2)	10	1.94	1.94
AgI	4-4	5-4	10	2.80	2.81
CdS	4-4	3½-3½	10	2.51	2.51
CdSe	4-4	4-4	10	2.66	2.63
InN	4-4	3(2)	10	2.15	2.13

A feature of Table 9 is the appearance of one of the normally electro-positive elements of group 2B, aluminum, in the role of a Division III element. Beryllium and magnesium also form ZnS type compounds, but like the lithium compounds previously mentioned they are irregular, probably for the same reason, and have not been included in the tabulation. The Division III behavior of these normally Division I elements is a result of the small size of the lower groups, which puts their Division I elements into the same positions with respect to the electronegative zero point as the Division III elements of the larger groups. This relation is indicated in the following tabulation, where the asterisks identify those elements that are normally in Division I:

DIVISION III

Be*	Mg*	Zn
B*	Al*	Ga
C	Si	Ge
N	P	As
O	S	Se
F	Cl	Br

None of the orientations thus far considered is applicable to compounds of the Division II elements. The normal orientation does not exist above a specific rotation of 5, as the higher value would put the relative rotation above the limiting value 10. The zinc oxide and zinc sulfide types of combination are electronegative structures, and the reverse orientation of the Division II elemental structures is not available for compounds with negative elements. The Division II elements therefore form their compounds on the basis of the magnetic orientation. This type of structure is theoretically available for any element, but its use is limited by probability considerations. It is utilized in many of the compounds of Divisions III and IV, especially in the higher rotational groups, but rarely appears in Division I combinations because of the very high probability of the normal orientation in this division.

Since the magnetic rotation is distributed over all three dimensions, its effective component is not altered by a change in position, and has the same value in the magnetic orientations as in the corresponding compounds based on the electric orientations. In order to establish the magnetic type of equilibrium, however, the axis of the negative electric rotation has to be parallel to that of one of the magnetic rotations, and it is therefore perpendicular to the axis of the positive electric rotation. Consequently, the latter takes no part in the normal inter-atomic force equilibrium, and it constitutes an additional orienting influence, the effects of which were discussed in Volume I. In these compounds of the magnetic type the displacement of the negative component ($-x$) is balanced by a numerically equal positive displacement (x). Thus the magnetic orientation is somewhat similar to the normal orientation. However, the magnetic rotation is opposite in vectorial direction to the electric rotation, and the resultant relative rotation effective in the dimension of combination is therefore one of the neutral values 10, 5, or a combination of these two, rather than the $2x$ of the normal orientation.

Compounds based on the magnetic orientation occur in a variety of crystal forms, the nature of which depends on the degree of force symmetry and the number of atoms of each kind in the equilibrium system. In some cases there is enough symmetry to make isometric structures of the NaCl , CaF_2 , and similar types possible. Other

crystals are asymmetric. A common arrangement for the binary compounds is the *nickel arsenide* structure, a hexagonal crystal in which the positive atoms occupy the face positions and the negative atoms are in the central positions, spaced alternately $\frac{1}{4}$ and $\frac{3}{4}$ along the c axis. Table 10 shows the inter-atomic distances calculated for some NiAs and NaCl type crystals of binary magnetic orientation compounds of Group 3A.

TABLE 10: *DISTANCES—BINARY MAGNETIC ORIENTATION COMPOUNDS*

Compound	Specific Rotation			Distance	
	Magnetic		Electric	Calc.	Obs.
NiAs (Hexagonal) Class—Group 3A					
VS	4-3	3½-3½	10	2.42	2.42
VSe	4-3	4-4	10	2.56	2.55
CrS	4-3	3½-3½	10	2.42	2.44
CrSe	4-3	4-4	10	2.56	2.54
CrSb	4-3	5-4½	10	2.73	2.74
CrTe	4-3	5-4½	10	2.73	2.77
MnAs	4-3	4-4	10	2.56	2.58
MnSb	4-3	5-4½	10	2.73	2.78
FeS	4-3	3½-3½	10	2.42	2.45
FeSe	4-3	4-4	10	2.56	2.55
FeSb	4-3	5-4	10	2.69	2.67
FeTe	3-4	5-4	10	2.59	2.61
CoS	3-4	3½-3½	10	2.32	2.33
CoSe	3-4	4-4	10	2.46	2.46
CoSb	3-4	5-4	10	2.59	2.58
CoTe	3-4	5-4	10	2.59	2.62
NiS	3½-3½	3½-3½	10	2.37	2.38
NiAs	3½-3½	4-3	10	2.42	2.43
NiTe	3½-3½	5-4	10	2.64	2.64
NaCl (Cubic) Class—Group 3A					
VN	4-3	3(2)	10	2.04	2.06
VO	4-3	3(2)	10	2.04	2.05
CrN	4-3	3(2)	10	2.04	2.07
MnO	3½-3½	3(2)	5-10	2.18	2.22
MnS	3½-3½	3½-3½	5-10	2.59	2.61
MnSe	3½-3½	4-4	5-10	2.75	2.72
FeO	3-4	3(2)	5-10	2.12	2.16
CoO	3-4	3(2)	5-10	2.12	2.12

TABLE 11: *DISTANCES-BINARY MAGNETIC ORIENTATION COMPOUNDS*

Compound	Specific Rotation		Distance	Calc.	Obs.
	Magnetic	Electric			
NaCl (Cubic) Class—Groups 4A and 4B					
CeN	5-4	3(2)	5-10	2.52	2.50
CeP	5-4	4-3	5-10	2.94	2.95
CeS	5-4	3½-3½	5-10	2.89	2.89*
CeAs	5-4	4-4	5-10	3.06	3.03
CeSb	5-4	5-4	5-10	3.22	3.20
CeBi	5-4	5-4	5-10	3.22	3.24
PrN	5-4	3(2)	5-10	3.52	2.58
PrP	5-4	4-3	5-10	2.94	2.93
PrAs	4½-4	4-4	5-10	2.98	3.00
PrSb	4½-4	5-4	5-10	3.14	3.17
NdN	5-4	3(2)	5-10	2.52	2.57
NdP	5-4	4-3	5-10	2.94	2.91
NdAs	4½-4	4-4	5-10	2.98	2.98
NdSb	4½-4	5-4	5-10	3.14	3.15
EuS	5-4	4-3	5-10	2.94	2.98
EuSe	5-4	4-4	5-10	3.06	3.08
EuTe	5-4	5-4½	5-10	3.26	3.28
GdN	5-4	3(2)	5-10	2.52	2.50*
YbSe	4½-4	4-4	5-10	2.98	2.93
YbTe	4½-4	5-4	5-10	3.14	3.17
ThS	4½-4½	3½-3½	5-10	2.85	2.84
ThP	4½-4½	4-3	5-10	2.91	2.91
UC	4½-4½	3(2)	5-10	2.47	2.50*
UN	4½-4½	3(2)	5-10	2.47	2.44*
UO	4½-4½	3(2)	5-10	2.47	2.46*
NpN	4½-4½	3(2)	5-10	2.47	2.45*
PuC	4½-4½	3(2)	5-10	2.47	2.46*
PuN	4½-4½	3(2)	5-10	2.47	2.45*
PuO	4½-4½	3(2)	5-10	2.47	2.48*
AmO	4½-4½	3(2)	5-10	2.47	2.48*

Almost all of the NiAs type compounds that have been examined in the course of this present work take the vibration one value of the specific electric rotation: 10. The magnetic orientation compounds with

the NaCl structure are quite evenly divided between the 10 rotation and the combination 5-10 in the 3A group, but utilize the 5-10 rotation almost exclusively in the higher groups. In order to show as wide a variety of the features of these magnetic type compounds as is possible in the limited amount of space that can be allotted to them, Table 10 has been restricted to Group 3A compounds, and the following Table 11 gives the data for a representative sample of the compounds of the rare earth elements (from Group 4A), together with a selection of compounds from Group 4B, in which the identical values of the inter-atomic distance in the combinations of the elements of this group with the Division IV elements of Group 2A are emphasized.

Thus far the calculation of equilibrium distances has been carried out by crystal types as a matter of convenience in identifying the effect of various atomic characteristics on the crystal form and dimensions. It is apparent from the points brought out in the discussion, however, that identification of the crystal type is not always essential to the determination of the inter-atomic distance. For example, let us consider the series of compounds NaBr, Na₂Se, and Na₃As. From the relations that have been established in the preceding pages we may conclude that these Division I compounds are formed on the basis of the normal orientation. We therefore apply the known value of the relative specific electric rotation of a normal orientation sodium compound, 4, and the known values of the normal specific magnetic rotations of sodium and the Group 3B elements, 3-3½ and 4-4 respectively, to equation 1-10, from which we ascertain that the most probable inter-atomic distance in all three compounds is 2.95, irrespective of the crystal structure. (Measured values are 2.97, 2.95, and 2.94 respectively.)

The possible inter-atomic distances in the more complex compounds can be calculated in a similar manner, without the necessity of analyzing the great variety of geometrical structures in which these compounds crystallize. The usefulness of this procedure in application to compounds in general is limited, at the present stage of the theoretical development, because we are not normally able to define the specific rotations from theoretical premises as definitely as in the foregoing illustration. It is of considerable value, however, in dealing with the lower electronegative elements, whose specific electric rotations are confined to the neutral values, and whose variability in the magnetic dimensions is only in

the number of inactive dimensions (that is, dimensions in which the specific rotation is 2). The elements involved are those of groups 1B and 2A; hydrogen, carbon, nitrogen, oxygen, and fluorine, together with boron, one of the normally electropositive elements of Group 2A. The other two positive elements of this group, lithium and beryllium, are also two-dimensional under most conditions, but they take the positive orientation, and have much greater inter-atomic distances.

Table 12 gives the theoretically possible inter-atomic distances of these lower group elements, with some examples of the measured values corresponding to the calculated distances.

TABLE 12: *DISTANCES-LOWER NEGATIVE ELEMENTS*

Specific Rotation				Distance			
		Magnetic	Elec.			n.u.	Å
	3(1)	3(1)	10			.241	.70
	3(1)	3(1½)	10			.317	.92
	3(1½)	3(1½)	10			.363	1.06
	3(1)	3(2)	10			.406	1.18
	3(1½)	3(2)	10			.445	1.30
	3(2)	3(2)	10			.483	1.41
	3(2)	3(2)	5-10			.528	1.54
Calc.	Comb.	Example	Obs.	Calc.	Comb.	Example	Obs.
.70	H-H	H ₂	.74	1.30	H-B	B ₂ H ₆	1.27
.92	H-F	HF	.92		C-O	CaCO ₃	1.29
	H-C	Benzene	.94		B-F	BF ₃	1.30
	H-O	Formic Acid	.95		C-N	Oxamide	1.31
1.06	H-N	Hydrazine	1.04		C-F	CF ₃ Cl	1.32
	H-C	Ethylene	1.06		C-C	Ethylene	1.34
	C-N	NaCN	1.09	1.41	C-C	Benzene	1.39
	N-N	N ₂	1.09		N-O	HNO ₃	1.41
	C-O	COS	1.10		C-C	Graphite	1.42
1.18	C-O	CO	1.15		C-N	dl-Alanine	1.42
	C-N	Cyanogen	1.16		C-O	Methyl ether	1.42
	H-B	B ₂ H ₆	1.17		C-F	CH ₃ F	1.42
	N-N	CuN ₃	1.17	1.54	C-C	Diamond	1.54
	N-O	N ₂ O	1.19		C-C	Propane	1.54
	C-C	Acetylene	1.20		B-C	B(CH ₃) ₂	1.56

The experimental results are not all in agreement with the theory. On the contrary, they are widely scattered. The measured C-C distances, for example, cover almost the entire range from 1.18, the minimum for this combination, to the maximum 1.54. However, the *basic* compounds of each class do agree with the theoretical values. The paraffin hydrocarbons, benzene, ethylene, and acetylene, have C-C distances approximating the theoretical 1.54, 1.41, 1.30, and 1.18 respectively. All C-H distances are close to the theoretical 0.92 and 1.06, and so on. It can reasonably be concluded, therefore, that the significant deviations from the theoretical values are due to special factors that apply to the less regular structures.

A detailed investigation of the reasons for these deviations is beyond the scope of this present work. However, there are two rather obvious causes that are worth mentioning. One is that forces exerted by adjacent atoms may modify the normal result of a two-atom interaction. An interesting point in this connection is that the effect, where it occurs, is inverse; that is, it increases the atomic separation, rather than decreasing it as might be expected. The natural reference system always progresses at unit speed, irrespective of the positions of the structures to which it applies, and consequently the inward force due to this progression always remains the same. Any interaction with a third atom introduces an additional rotational (outward) force, and therefore moves the point of equilibrium outward. This is illustrated in the measured distances in the polynuclear derivatives of benzene. The lowest C-C distances in these compounds, 1.38 and 1.39, are found along the outer edges of the molecular structures, while the corresponding distances in the interiors of the compounds, where the influence of adjoining atoms is at a maximum, characteristically range from 1.41 to 1.43.

Another reason for discrepancies is that in many instances the measurement and the theoretical calculation do not apply to the same quantity. The calculation gives us the distance between structural units, whereas the measurements apply to the distances between specific atoms. Where the atoms are the structural units, as in the compounds of the NaCl class, or where the inter-group distance is the same as the inter-atomic distance, as in the normal paraffins, there is no problem, but exact agreement cannot be expected otherwise. Again we can use benzene as an example. The C-C distance in benzene is generally

reported as 1.39, whereas the corresponding theoretical distance, as indicated in Table 12, is 1.41. But, according to the theory, benzene is not a ring of carbon atoms with hydrogen atoms attached; it is a ring of CH neutral groups, and the 1.41 neutral value applies to the distance between these neutral groups, the structural units of the atom. Since the hydrogen atoms are known to be outside the carbon atoms, if these atoms are coplanar it follows that the distance between the effective centers of the CH groups must be somewhat greater than the distance between the carbon atoms of these groups. The 1.39 measurement between the carbon atoms is therefore entirely consistent with the theoretical distance calculations.

The same kind of a deviation from the results of the (apparent) direct interaction between two individual atoms occurs on a larger scale where there is a group of atoms that is acting structurally as a radical. Many of the properties of molecules composed in part, or entirely, of radicals or neutral groups are not determined directly by the characteristics of the atoms, but by the characteristics of the groups. The NH_4 radical, for example, has the same specific rotations, when acting as a group, as the rubidium atom, and it can be substituted in the NaCl type crystals of the rubidium halides without altering the volume. Consequently, the inter-atomic distances have no direct significance in compounds containing these groups. It is theoretically feasible to locate the effective centers of the various groups, and to measure the inter-group distances that correspond to those calculated from theory, but this task has not yet been undertaken, and it will not be possible at this time to present a comparison between theoretical and experimental distances in compounds containing radicals comparable to the comparisons in Tables 1 to 12.

Some preliminary results have been made, however, on the relation between the theoretical distances and the *density* in complex compounds. There are a number of factors, not yet investigated in detail, that have some influence on the density of solid matter, and for that reason the conclusions thus far derived from theory are somewhat tentative, and the correlations between theory and observation are only approximate. Nevertheless, certain aspects of these tentative results are significant, and are of enough interest to justify giving them some attention.

If we divide the molecular mass, in terms of atomic weight units, by the density, we arrive at the molecular volume in terms of the units entering into the density measurement. For present purposes it will be convenient to convert this quantity to natural units of volume. The applicable conversion factor is the cube of the time region unit of distance divided by the mass unit atomic weight. In the cgs system of units it has the numerical value 14.908.

In Table 13 the average volumes per volumetric group of a representative number of inorganic compounds containing radicals (V), as calculated from the measured densities, are compared with the cubes of the inter-group distances (S_0^3), as calculated on the theoretical basis previously described.

TABLE 13: *MOLECULAR VOLUME*

	m	d	n	V	s_0^3	c	ab_1	ab_2
NaNO ₃	85.01	2.261	2	1.261	1.241	4	3-3	4-5
KNO ₃	101.10	2.109	2	1.608	1.565	4	4-3	4-5
Ca(NO ₃) ₂	164.10	2.36	3	1.554	1.565	4	4-3	4-5
RbNO ₃	147.49	3.11	2	1.590	1.631	4	4-4	4-4
Sr(NO ₃) ₂	211.65	2.986	3	1.585	1.631	4	4-4	4-4
CsNO ₃	194.92	3.685	2	1.774	1.825	4	4½-4½	4-4
Na ₂ CO ₃	106.00	2.509	3	0.944	0.970	4	3-3	3½-3½
MgCO ₃	84.33	3.037	2	0.931	0.970	4	3-3	3½-3½
K ₂ CO ₃	138.20	2.428	3	1.272	1.222	4	4-3	3½-3½
CaCO ₃	100.09	2.711	2	1.238	1.222	4	4-3	3½-3½
BaCO ₃	197.37	4.43	2	1.494	1.532	4	4½-4½	3½-3½
FeCO ₃	115.96	3.8	2	1.022	0.976	4	4-3	3½-3½
CoCO ₃	118.95	4.13	2	0.966	0.976	4	4-3	3½-3½
Cu ₂ CO ₃	187.09	4.40	3	0.950	0.976	4	4-3	3½-3½
ZnCO ₃	125.39	4.44	2	0.947	0.976	4	4-3	3½-3½
Ag ₂ CO ₃	275.77	6.077	3	1.015	1.096	4	4-4	3½-3½

The specific electric rotation (c) for the compounds with the normal orientation is 4, as in the valence one binary compounds. Those with the magnetic orientation take the neutral value 5. The applicable specific magnetic rotations for the positive component

and the negative radical are shown in the columns headed ab_1 and ab_2 respectively. Columns 2, 3, and 4 give the molecular mass (m), the density of the solid compound (d), and the number of volumetric units in the molecule (n). Here, again, as in the earlier tables, the calculated and empirical values are not exactly comparable, as the measured values of the densities have been used directly, rather than being projected back to zero temperature, a refinement that would be required for accuracy, but is not justified at this early stage of the investigation.

In this table there are five pairs of compounds, such as $\text{Ca}(\text{NO}_3)_2$ and KNO_3 in which the inter-group distances are the same, and the only difference between the pairs, so far as the volumetric factors are concerned, is in the number of structural groups. Because of the uncertainties involved in the measured densities, it is difficult to reach firm conclusions on the basis of each pair considered individually, but the average volume per group, calculated from the density, in the five two-group structures is 1.267, whereas in the five three-group structures the average is 1.261. It is evident from this that the volumetric equality of the group and the independent atom which we noted in the case of the NH_4 radical is a general proposition, in this class of compounds at least. This is a point that will have a special significance when we take up consideration of the liquid volume relations.

In closing the discussion in this chapter it is appropriate to reiterate that the values of the inter-atomic and inter-group distance derived from theory apply to the separations as they would exist if the equilibrium were reached at zero temperature and zero pressure. In the next two chapters we will consider how these distances are modified when the solid structure is subjected to finite pressures and temperatures.

CHAPTER 4

Compressibility

One of the simplest physical phenomena is *compression*, the response of the time region equilibrium to external forces impressed upon it. With the benefit of the information developed earlier, we are now in a position to begin an examination of the compression of solids, disregarding for the present the question of the origin of the external forces. For this purpose we introduce the concept of *pressure*, which is defined as force per unit area.

$$P = \frac{F}{s^2} \quad (4-1)$$

In many cases it will be convenient to deal with pressure on a volume basis rather than on an area basis. We therefore multiply both force and area by distance, s , which gives us the alternative equation:

$$P = \frac{Fs}{s^3} = \frac{E}{V} \quad (4-2)$$

In the region outside unit distance, where the atoms or molecules of matter are independent, the total energy of an aggregate can thus be expressed in terms of pressure and volume as

$$E = PV \quad (4-3)$$

As we will find in the next chapter when we begin consideration of thermal motion, a condition of constant temperature is a condition of constant energy, other things being equal. Equation 4-3 thus tells us that for an aggregate in which the cohesive forces between the atoms or molecules are negligible, an *ideal gas*, the volume at constant temperature is inversely proportional to the pressure. This is *Boyle's law*, one of the well-established relations of physics.

For application to the time region in which the solid equilibrium is located, the second power of the volume must be substituted for

the first power, in accordance with the general inter-regional relation previously established. The time region equivalent of Boyle's Law is therefore

$$PV^2 = k \quad (4-4)$$

In terms of volume this becomes

$$V = \frac{k}{P^{1/2}} \quad (4-5)$$

This equation tells us that at constant temperature the volume of a solid is inversely proportional to the square root of the pressure. The pressure represented by the symbol P in this equation is, of course, the total effective pressure; that is, the pressure equivalent of *all* of the forces acting in opposition to the rotational forces of the atom. The force due to the progression of the natural reference system opposes the rotational forces, and acts in parallel with the external compressive forces, but it has the same magnitude regardless of whether or not any such external forces are present. It therefore exerts what we may call an *internal pressure*, an already existing level of pressure to which an external pressure becomes an addition. In order to conform to established usage and to avoid confusion, the symbol P will hereafter refer to the external pressure only, the total pressure being expressed as $P_0 + P$. On this basis, equation 4-5 may be restated as

$$V = \frac{k}{(P_0 + P)^{1/2}} \quad (4-6)$$

Compression is normally expressed in terms of relative rather than absolute volumes, the reference volume being the volume at zero external pressure, where equation 4-6 has the form

$$V_0 = \frac{k}{P_0^{1/2}} \quad (4-7)$$

Dividing equation 4-6 by equation 4-7, and rearranging, we obtain

$$\frac{V}{V_0} = \frac{P_0^{1/2}}{(P_0 + P)^{1/2}} \quad (4-8)$$

As this equation brings out, the internal pressure, P_0 , is the key factor in the compression of solids. Inasmuch as this pressure is a result of

the progression of the natural reference system which, in the time region, is carrying the atoms inward in opposition to their rotational forces (gravitation), the inward force acts only on two dimensions (an area), and the magnitude of the pressure therefore depends on the orientation of the atom with respect to the line of the progression. As indicated in connection with the derivation of the inter-regional ratio, there are 156.44 possible positions of a displacement unit in the time region, of which a fraction az represents the area subjected to pressure, a and z being the effective displacements in the active dimensions. The letter symbols a , b , and c , are used as indicated in Chapter 10, Volume I. The displacement z is either the electric displacement c or the second magnetic displacement b , depending on the orientation of the atom.

From the principle of equivalence of natural units it follows that each natural unit of pressure exerts one natural unit of force per unit of cross-sectional area per effective rotational unit in the third dimension of the equivalent space. However, the pressure is measured in the units applicable to the effect of external pressure. The forces involved in this pressure are distributed to the three spatial dimensions and to the two directions in each dimension. Only those in one direction of one dimension—one sixth of the total—are effective against the time region structure. Applying this $\frac{1}{6}$ factor to the ratio $az/156.444$, we have for the internal pressure per rotational unit at unit volume,

$$P_0 = \frac{az}{938.67} \quad (4-9)$$

This expression may now be generalized to apply to y rotational units and V units of volume, as follows:

$$P_0 = \frac{azy}{938.67 V} \quad (4-10)$$

The force exerted by the progression of the natural reference system is independent of the geometrical arrangement of the atoms, and the volume term in equation 4-10 refers to what we may call the three-dimensional atomic space, the cube of the inter-atomic distance, rather than to the geometrical volume. We will therefore replace V by s_0^3 . This gives us the internal pressure equation in final form:

$$P_0 = \frac{azy}{938.67 s_0^3} \quad (4-11)$$

The value derived from this equation is the magnitude of the internal pressure in terms of natural units. To obtain the pressure in terms of any conventional system of units it is only necessary to apply a numerical coefficient equal to the value of the natural unit of pressure in that system. This natural unit was evaluated in Volume I as 1.575×10^{13} dynes/cm². The corresponding values in the systems of units used in the reports of the experiments with which comparisons will be made in this chapter are:

$$\begin{aligned} &1.554 \times 10^7 \text{ atm} \\ &1.606 \times 10^7 \text{ kg/cm}^2 \\ &1.575 \times 10^7 \text{ megabars} \end{aligned}$$

In terms of the units used by P. W. Bridgman, the pioneer investigator in the field, in most of his work, equation 4-11 takes the form

$$P_0 = \frac{17109 azy}{s_0^3} \frac{\text{kg}}{\text{cm}^2} \quad (4-12)$$

The internal pressure thus calculated for any specific substance is not usually constant through the entire external pressure range. At low total pressures, the orientation of the atom with respect to the line of progression of the natural reference system is determined by the thermal forces which, as we will see later, favor the minimum values of the effective cross-sectional area. In the low range of total pressures, therefore, the cross-section is as small as the rotational displacements of the atom will permit. In accordance with *Le Chatelier's Principle*, a higher pressure, either internal or external, applied against the equilibrium system causes the orientation to shift, in one or more steps, toward higher displacement values. At extreme pressures the compressive force is exerted against the maximum cross-section: 4 magnetic units in one dimension and 8 electric units in another. Similarly, only one of the magnetic rotational units in the atom participates in the radial component y of the resistance to compression at the low pressures, but further application of pressure extends the participation to additional rotational units, and at extreme pressures all of the rotational units in the atom are involved. The limiting value of y is therefore the total number of such units. The exact sequence in

which these two kinds of factors increase in the intermediate pressure range has not yet been determined, but for present purposes a resolution of this issue is not necessary, as the effect of any specific amount of increase is the same in both cases.

Helium and neon, the first two of the inert gases, the elements that have no effective rotation in the electric dimension, take the absolute minimum compression factors: one rotating unit with one effective unit of displacement in each of the two effective dimensions. The *azy* factors for these elements can be expressed as 1-1-1. In this notation, which we will use for convenience in the subsequent discussion, the numerical values of the compression factors are given in the same order as in the equations. It should be noted that the absolute minimum compression, that applicable to the elements of least displacement, is *explicitly* defined by the factors 1-1-1. The value of the factor *a* increases in the higher members of the inert gas series because of their greater magnetic displacement.

Because of their negative displacement in the electric dimension, which, in this context, is equivalent to the zero displacement of the inert gases, the electronegative elements follow the inert gas pattern, taking the minimum 1-1-1 factors in the lowest members of the lowest rotational groups, and values that are higher, but still generally well below those of the corresponding electropositive elements, as the displacement increases in either or both of the atomic rotations. None of the elements of the electronegative divisions below electric displacement 7 has the 4-8 *az* factors initially, although they are capable of these high levels, and can eventually reach them under appropriate conditions.

All of the electropositive elements studied by Bridgman have the full 4 units in one dimension; that is, $a=4$. The value of *z* for the alkali metals is equal to the electric displacement, one unit, and since *y* takes the minimum value under low pressure conditions, the compression factors for these elements are 4-1-1. The displacement 2 elements (calcium, etc.) take the intermediate values 4-2-1 or 4-3-1. The greater displacements of the elements that follow have a double effect. They increase the internal pressure directly by enlarging the effective cross-section, and this higher internal pressure then has the same effect as a greater external pressure in causing a further

increase in the compression factors. Most of these elements therefore utilize the full displacements of the active cross-section dimensions from the start of compression; that is, 4-4-1 ($az=ab$, two magnetic dimensions) in some of the lower group elements and the transition elements of Group 4A, and 4-8-1, or 4-8-n ($az=ac$, one magnetic and one electric dimension) in the others.

The factors that determine the internal pressures of the compounds that have been investigated thus far fall mainly in the intermediate range, between 4-1-1 and 4-4-1. NaCl, for instance, has 4-2-1 initially, and shifts to 4-3-1 in the pressure range between 30 and 50 M kg/cm². AgCl has 4-3-1 initially, and carries these factors up to a transition point near Bridgman's pressure limit of 100 M kg/cm². CaF₂ has the factors 4-4-1 from the start of compression. The initial values of the internal pressure of most of the inorganic compounds examined in this investigation are based on one or another of these three patterns. Those of the organic compounds are mainly 4-1-1, 4-2-1, or an intermediate value 4-1½-1.

Compression is ordinarily measured in terms of relative volume, and most of the discussion in this chapter will deal with the subject on this basis, but for some purposes we will be interested in the *compressibility*, the rate of change of volume under pressure. This rate is obtained by differentiating equation 4-8.

$$\frac{1}{V_0} \frac{dV}{dP} = \frac{P_0^{1/2}}{2(P_0 + P)^{3/2}} \quad (4-13)$$

The compressibility at P_0 , the *initial compressibility*, is of particular interest. For all practical purposes it is the same as the compressibility at one atmosphere, this pressure being only a small fraction of the internal pressure P_0 . The initial compressibility may be obtained from equation 4-13 by letting P equal zero. The result is

$$\frac{1}{V_0} \frac{dV}{dP} (P = 0) = \frac{1}{2P_0} \quad (4-14)$$

Since the initial compressibility is a quantity that can be measured, its simple and direct relation to the internal pressure provides a significant confirmation of the physical reality of that theoretical property of matter. Initial compressibility factors derived theoretically for those elements on which consistent compressibility

data are available for comparison, the internal pressures calculated from these factors, and the initial compressibilities corresponding to the calculated internal pressures are listed in Table 14, together with measured values of the initial compressibility at room temperature. Two sets of experimental values are given, one from Bridgman and one from a more recent compilation. Values of s_0^3 , except those marked with asterisks, are computed from the inter-atomic distances (s_0) in the tables of Chapter 2. Where the structure is anisotropic the s_0^3 value shown is the product of one of the distances given in the earlier tabulations by the square of the other. The reason for the occurrence of the indicated deviations from the Chapter 2 values will be explained later.

TABLE 14: *INITIAL COMPRESSIBILITY*

	S_0^3	Comp. factors			P_0 (M kg/cm ²)	Initial compressibility $\times 10^6$		
		a	z	y		Calc.	Obs. ³	Obs. ⁴
Li	1.151	4	1	1	59.5	8.42	8.41	8.46
Be	0.482	4	4	1	568	0.88	0.87	0.98
C(dia.)	0.147	4	6	1	2793	0.18	0.18	0.18
Na	2.048	4	1	1	33.4	14.97	15.1	14.42
Mg	1.291	4	4	1	212	2.36	2.86	2.77
Al	0.915	4	5	1	374	1.34	1.30	1.36
Si	0.497	4	4	1	551	0.91	0.31	0.99
K	3.659	4	1	1	18.7	26.74	31.0	30.4
Ca	2.588	4	3	1	79.3	6.31	5.51	6.45
Ti	1.033	4	8	1	530	0.94	0.77	0.93
V	0.729	4	8	1	751	0.67	0.59	0.61
Cr	0.603	4	8	1	908	0.55	0.50	0.52
Mn	0.705	4	8	1	777	0.64	0.76	1.65
Fe	0.603	4	8	1	908	0.55	0.57	0.58
Co	0.603*	4	8	1	908	0.55	0.52	0.51
Ni	0.603*	4	8	1	908	0.55	0.50	0.53
Cu	0.652	4	6	1	630	0.79	0.70	0.72
Zn	0.903	4	4	1	303	1.65	1.64	1.64
Ge	0.603	4	4	1	454	1.10	1.33	1.27
Rb	4.616	4	1	1	14.8	33.78	38.7	31.4
Sr	3.268	4	3	1	62.8	7.96	7.9	8.46
Zr	1.306	4	6	1½	472	1.06	1.06	1.18

	S_0^3	Comp. factors			P_0 (M kg/cm ²)	Initial compressibility $\times 10^6$		
		a	z	y		Calc.	Obs. ³	Obs. ⁴
Mo	0.764*	4	8	2	1433	0.35	0.34	0.36
Ru	0.764*	4	8	2	1433	0.35	0.34	0.31
Rh	0.764	4	8	2	1433	0.35	0.36	0.36
Pd	0.823	4	8	1½	998	0.50	0.51	0.54
Ag	0.956	4	8	1	573	0.87	0.96	0.97
Cd	1.118	4	4	1	245	2.04	1.89	2.10
In	1.165*	4	4	1	235	2.13		2.38
Sn	0.913*	4	4	1	300	1.67	1.64	0.80
Sb	1.325*	4	4	1	209	2.42	2.32	2.56
Cs	5.774	4	1	1	11.9	42.0	59.0	49.1
Ba	2.686*	4	2	1	51.0	9.80		9.78
La	2.044	4	4	1	134	3.73	3.39	4.04
Ce	1.893	4	4	1	145	3.45	3.45	4.10
Pr	1.758*	4	4	1	156	3.21		3.21
Nd	1.758*	4	4	1	156	3.21		3.00
Sm	1.758*	4	4	1	156	3.21		3.34
Gd	1.346*	4	4	1	203	2.46		2.56
Dy	1.346	4	4	1	203	2.46		2.55
Ho	1.346*	4	4	1	203	2.46		2.47
Er	1.346*	4	4	1	203	2.46		2.38
Tm	1.346*	4	4	1	203	2.46		2.47
Yb	2.167*	4	2	1	63.2	7.92		7.38
Lu	1.346*	4	4	1	203	2.46		2.38
Ta	1.027*	4	8	2	1066	0.47	0.47	0.49
W	0.953*	4	8	3	1723	0.29	0.28	0.30
Ir	0.823	4	8	3	1996	0.25		0.28
Pt	0.823	4	8	2	1330	0.38	0.35	0.35
Au	0.953	4	8	1½	862	0.58	0.56	0.57
Tl	1.631	4	4	1	168	2.98	3.31	2.74
Pb	1.249*	4	4	1	219	2.25	2.29	2.29
Bi	1.249	4	3	1	164	3.05	2.71	3.11
Th	1.758	4	8	1	311	1.61		1.81
U	0.984	4	8	1	556	0.90	0.94	0.99

In most cases the difference between the calculated and measured compressibilities is within the probable experimental error. Substantial deviations from the calculated values are to be expected in the case of low melting point elements such as the alkali metals, unless corrections have been applied to the empirical data, as there is an additional component in the initial volume of such substances. Elsewhere, the differences between the calculated compressibilities and either of the two sets of experimental values are, on the average, no greater than the differences between the experimental results.

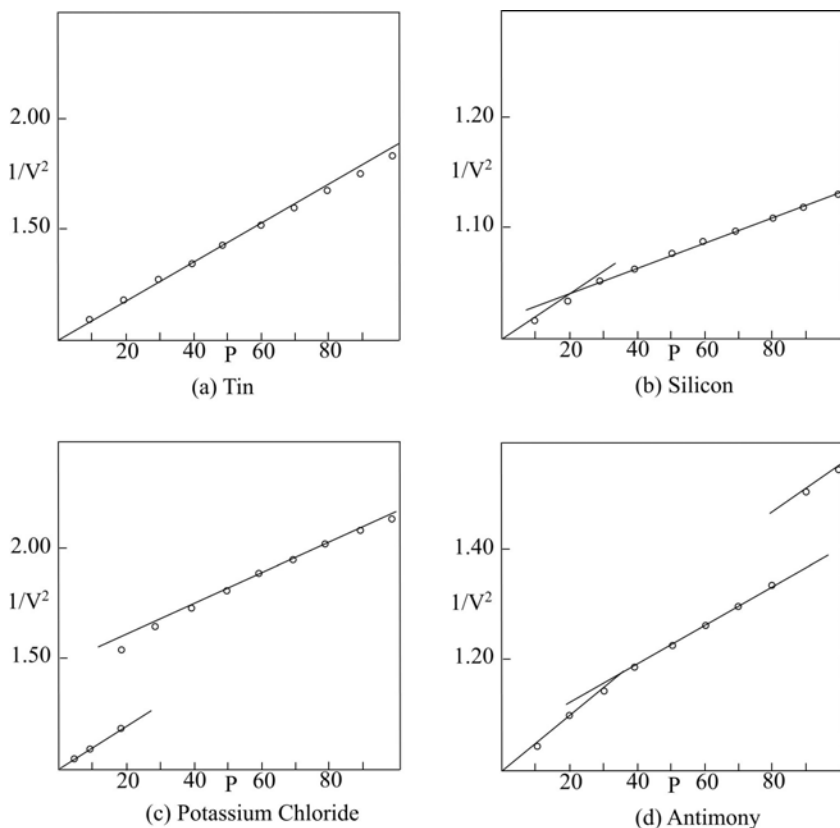
As indicated in the foregoing paragraphs, the general compression pattern can be described in this manner: Each element has an initial rate of compression based on the lowest set of compression factors permitted by the specific rotations of the atoms of that element. As the pressure increases, the volume decreases in proportion to the square root of the total pressure up to some specific pressure level, where one or more of the compression factors changes to a higher value, increasing the internal pressure proportionately. The compression then continues in proportion to the square root of the higher pressure. This process is repeated at successively higher pressure levels until the maximum compression factors for the element are reached.

Because of the nature of this compression pattern, a convenient method of analyzing the experimental values of the volume of various substances under compression can be made available by expressing equation 4-8 in the form

$$\left(\frac{V_0}{V}\right)^2 = 1 + \frac{P}{P_0} \quad (4-15)$$

According to this equation, if we plot the reciprocals of the squares of the relative volumes against the corresponding total pressure ratios we should obtain a straight line intersecting the zero pressure ordinate at the reference volume 1.00. The slope of the line is determined by the magnitude of the internal pressure, P_0 . Fig.1(a) is a curve of this kind for the element tin, based on Bridgman's experimental values.

Where there is a transition to a higher set of compression factors within the experimental range, and the magnitude of P_0 changes, the volumes diverge from the original line and follow a second straight line, the slope of which is determined by the new compression factors.

FIGURE 1: *COMPRESSION PATTERNS*

In preparing curves of this kind for the other elements investigated by Bridgman, we find that about two-thirds of them actually do conform to a single straight line up to the 30,000 kg/cm² pressure limit of his earlier work. His studies of the less compressible substances, such as the higher elements of the electropositive divisions, were not carried beyond this level, but he measured the compression up to 100,000 kg/cm² on many other elements, and most of them were found to undergo a transition in which the effective internal pressure increases without any volume discontinuity. The compression curve for such a substance consists of two straight line segments connected by a smooth transition curve, as in Fig. 1(b), which represents Bridgman's values for silicon.

In addition to the changes of this type, commonly called second order transitions, some solid substances undergo first order trans-

itions in which there is a modification of the crystal structure and a volume discontinuity at the transition point. The effective internal pressure generally changes during a transition of this kind, and the resulting volumetric pattern is similar to that of KCl, Fig.1(c). With the exception of some values which are rather erratic and of questionable validity, all of Bridgman's results follow one of these three patterns or some combination of them. The antimony curve, Fig.1(d), illustrates one of the combination patterns. Here a second order transition between 30,000 and 40,000 kg/cm² is followed by a first order transition at a higher pressure. The numerical values corresponding to these curves are given in the tables that follow.

The experimental second order curves are smooth and regular, indicating that the transition process takes place freely when the appropriate pressure is reached. The first order transitions, on the other hand, show considerable irregularity, and the experimental results suggest that in many substances the structural changes at the transition points are subject to a variable amount of delay due to internal conditions in the solid aggregate. In these substances the transition does not take place at a well-defined pressure, but somewhere within a relatively broad transition zone, and the exact course of the transition may vary considerably between one series of measurements and another. Furthermore, there are many substances which appear to experience similar delays in achieving volumetric equilibrium even where no transitions take place. The compression curves suggest that a number of the reported transitions are actually volume adjustments which merely reflect delayed response to the pressure applied earlier. For example, in the barium curve based on Bridgman's results there are presumably two transitions, one between 20,000 and 25,000 kg/cm², and the other between 60,000 and 70,000 kg/cm². Yet the experimental volumes at 60,000 and 100,000 kg/cm² are both very close to the values calculated on the basis of a single straight line relation. It is quite probable, therefore, that this element actually follows one linear relation at least up to the vicinity of 100,000 kg/cm².

The deviations from the theoretical curves that are found in the experimental volumes of substances with relatively high melting points are generally within the experimental error range, and those

larger deviations that do make their appearance can, in most cases, be explained on the foregoing basis. The compression curves for substances with low melting points show systematic deviations from linearity at the lower pressures, but this is a normal pattern of behavior resulting from the proximity of the change of state. As will be brought out in detail in our examination of the liquid state, the physical state of matter is basically a property of the individual atom or molecule. The state of the aggregate merely reflects the state of the majority of its individual constituents. Consequently, a solid aggregate at any temperature near the melting point contains a specific proportion of liquid molecules. Since the volume of a liquid molecule differs from that of a solid molecule, the volume of the aggregate is modified accordingly. The amount of the volume deviation in each case can be calculated by methods that will be described in the subsequent discussion of the liquid volume relations.

TABLE 15: *RELATIVE VOLUMES UNDER COMPRESSION*

Pressure (M kg/cm ²)	Calc.	Obs.	Calc.	Obs.	Calc.	Obs.	Calc.	Obs.
	Zn 4-4-1		Zr 4-6-1½		In 4-4-1		Sn 4-4-1	
0	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	.992	.992	.995	.995	.988	.988	.992	.991
10	.984	.984	.990	.989	.980	.980	.984	.982
15	.976	.977	.985	.983	.970	.967	.976	.975
20	.969	.969	.980	.978	.960	.955	.968	.966
25	.961	.964	.975	.973	.951	.948	.961	.960
30	.954	.954	.970	.969	.942	.936	.954	.951
35	.947		.965	.964	.933	.932	.947	
40	.940	.939	.960	.960	.925	.919	.940	.936
50	.927	.925	.951	.946	.909	.903	.926	.923
60	.914	.912	.942	.937	.893	.888	.913	.909
70	.902	.900	.933	.929	.978	.874	.901	.897
80	.890	.889	.925	.922	.864	.860	.889	.886
90	.879	.878	.917	.916	.851	.847	.878	.875
100	.868	.868	.909	.910	.838	.835	.867	.864

Table 15 compares the results of the application of equation 4-8 with Bridgman's measurements on some of the elements that maintain the same internal pressure all the way up to his pressure limit of 100,000 kg/cm². In many cases he made several series of measurements on the same element. Most of these results agree within about 0.003, and it does not appear that listing all of the individual values in the tabulations would serve any useful purpose. The values given in Table 15, and in the similar tables that follow, are those obtained in experiments that were carried to the full 100,000 kg/cm² pressure level. Where the high pressure measurements were started at some elevated pressure, or where the measurement interval was greater than usual, the gaps have been filled from the results of other Bridgman experiments.

Table 16 extends the volume comparisons to representative elements of the classes that are subject to transitions within the experimental range of pressures. Transitions reported by the investigator or indicated by the theoretical calculations are shown by horizontal lines in the appropriate columns. In these tabulations the position of the upper branch of each curve has been fixed by using the experimental volume at a selected pressure in the straight line segment above the transition (identified by the symbol R) as a reference point. Thus the *slope* of this upper branch of the curve is determined theoretically, but its *position* relative to the $1/\sqrt{v}$ scale is empirical. Some work has been done toward extension of the theoretical development to a determination of the exact position of the upper section of each curve, but this project is not far enough advanced to justify any discussion of it at this time.

Compressibility patterns of compounds are theoretically identical with those of the elements, and this theoretical conclusion is confirmed by compression data for a representative group of inorganic compounds presented in Table 17. As might be expected for the less uniform composition, transitions are somewhat more common in the compounds, but otherwise there is no difference in the compression curves. The curve for KCl, shown graphically in Fig. 1 and by numerical values in Table 17, is of special interest because it includes a sharp first order transition in which there is a substantial decrease in the basic volume while the compression factors remain unchanged. The magnitude of the volume reduction that takes place indicates that there is a reorientation of the atomic rotations in which the neutral specific electric rotation 5 is substituted for the normal rotation 4 as the effective relative value.

TABLE 16: *RELATIVE VOLUMES UNDER COMPRESSION*

Pressure (M kg/cm ²)	Calc.	Obs.	Calc.	Obs.	Calc.	Obs.	Calc.	Obs.
	Al		Si		Ca		Sb	
	4-5-1		4-4-1		4-3-1		4-4-1	
	4-8-1		4-8-1		4-4-1		4-4-1½	
0	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	.993	.993	.996	.995	.970	.969	.988	.987
10	.987	.987	.991	.990	.943	.942	.977	.975
15	.981	.981	.987	.986	<u>.917</u>	.918	.966	.964
20	.974	.975	.982	.981	.895	.897	.955	.954
25	<u>.968</u>	.969	.978	.978	.878	.878	.945	.944
30	.964	.964	.974	.974	.862	.861	.935	.934
35					.847	.845	.925	.925
40	.957	.958	<u>.966</u>	.968	.832	.832	<u>.916</u>	.917
50	.949	.951	.960	.962	.805 R	.805	.899	.899
60	.942	.944	.956	.957	.780	.780	.888	.886
70	.935	.937	.952	.952	.758	.748	.875	.875
80	.928	.929	.948	.948	.737	.732	.864 R	<u>.864</u>
90	.922	.922	.944	.944	.718	.716		.815
100	.915 R	.915	.940 R	.940	.701	.702		.803

	Al		Si		Ca		Sb	
	4-5-1		4-4-1		4-3-1		4-4-1	
	4-8-1		4-8-1		4-4-1		4-4-1½	
0	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
5	.955	.955	.982	.981	.984	.983	.996	.955
10	.915	.914	.965	.963	.970	.967	.991	.990
15	.880	.879	.949	.947	.955	.953	.987	.986
20	.848	<u>.841</u>	.933	<u>.933</u>	.942	.940	.983	.981
25	.820	.814	.918	.917	.929	.927	.979	.978
30	.794	.789	.904	.905	.916	.915	.975	.973
35	.771	.770	.891	.893	.904	.904	.971	.971
40	.750	.747	<u>.878</u>	.881	<u>.893</u>	.893	<u>.967</u>	.966
50	.712	.712	.858	.863	.878	.878	.960	.960
60	.679	.682	.845	.846	.863	.863	.956	.955
70	.950	.639	.833	.832	.849 R	.849	.952	.951
80	.625	.618	.821	.819	.835	.836	.949	.947
90	.603	.598	.809	.808	.822	.836	.945	.944
100	.582	.580	.798 R	.798	.810	.811	.941 R	.941

TABLE 17: *RELATIVE VOLUMES UNDER COMPRESSION*

Pressure (M kg/cm ²)	Calc.	Obs.	Calc.	Obs.	Calc.	Obs.	Calc.	Obs.
	NaCl		NaI		KCl		ZnS	
	4-2-1		4-2-1		4-2-1		4-4-1	
	4-2-1½		4-2-1½		4-2-1½		4-4-1½	
0	.994	1.000	.987	1.000	.994	1.000	.995	1.000
5	.979	.982	.964	.970	.973	.974	.991	.994
10	.964	.966	.942	.944	.953	.952	.986	.988
15	.950	.951	.922	.922	.934	.933	.982	.982
20	.937 R	.937	.903 R	.902	.916 R	<u>.916</u>	.977 R	.977
					.803 R	.803		
25	.924	.924	.885	.886	.791	.789	.973	.972
30	.912	.912	.868	.871	.779	.778	.969	.967
35	.900	.901	<u>.853</u>	.858	.768	.768	.964	.963
40	.889	.892	.840	.840	<u>.757</u>	.758	.960	.961
50	.867	.865	.819	.816	.741	.742	<u>.952</u>	.954
60	.847	.848	.799	.795	.727	.723	.945	.947
70	<u>.829</u>	.832	.781	.777	.714	.710	.940	.940
80	.815	.817	.765	.761	.702	.698	.934	.934
90	.802	.803	.749	.747	.690	.688	.929	.929
100	.790 R	.790	.734 R	.734	.679 R	.679	.924 R	.924
	AgCl		CsBr		NH ₄ Cl		KNO ₃	
	4-3-1		4-3-1		4-2-1		4-3-1	
			4-4-1		4-4-1		4-3-2	
0	1.000	1.000	.984	1.000	1.000	1.000	<u>.894</u>	1.000
5	.990	.989	.962	.971	.974	.973	.878	.882
10	.980	.979	.942	.947	.950	.951	.862	.862
15	.971	.969	.923	.925	<u>.928</u>	.933	.847	.846
20	.961	.960	.905 R	.905	.910	.918	.833	.831
25	.952	.952	.888	.888	.900	.905	.820 R	.820
30	.944	.942	.871	.870	.889	.891	.807	.804
35	.935	.937	.856	.859	.879	.883		
40	.927	.926	<u>.842</u>	.840	<u>.869</u>	.867	<u>.783</u>	.781
50	.911	.910	.815	.814	.851	.846	<u>.761</u>	.762
60	.895	.896	.790	.792	.833	.828	.744	.745
70	.881	.883	.777	.773	.817	.812	.733	.732
80	.867	.871	.760	.757	.801	.798	.723	.720
90	.854	.860	.743	.742	.787	.785	.712	.711
100	.841	.835	.728 R	.728	.773 R	.773	.703 R	.703

The theoretical volumes beyond the transition point, as shown in the table, are based on the small atomic volume corresponding to the higher rotation. Up to 20,000 kg/cm² the volume follows the curve corresponding to compression factors 4-2-1 and $s_0^3 = 1.222$, which produce an internal pressure of 112.7 M kg/cm². At the transition point the basic volume (s_0^3) drops to 0.976, increasing the internal pressure to 141.1 M kg/cm². The compression then continues on this basis up to the vicinity of 45,000 kg/cm², where the compression factors change from 4-2-1 to 4-3-1, and the internal pressure rises accordingly.

As in the compression of the elements, the theoretical calculations do not always confirm the transitions reported by the experimenters. On the other hand, these calculations show that a large proportion of the compounds, including six of the eight in Table 17, undergo either a transition or some other process in which they eliminate a volume component in the pressure range below 5000 kg/cm². The effect on the compression curve is to cause the linear segment of the curve to intersect the zero pressure ordinate at a volume below 1.000. The origin of these volume adjustments is still uncertain. The occurrence of a number of observable first order transitions at relatively low pressures suggests that some early second order transitions may also be taking place. But it is also possible that voids in the structure may be eliminated in the early stages of compression, or that there are geometrical readjustments.

The structural characteristics of the organic compounds make them particularly susceptible to such geometrical readjustments. Because of their low melting points, their volumes under low pressure also include the additional component that exists near the change of state. It appears, however, that in a wide range of compounds elimination of these extra volume components is substantially complete at some pressure well below the 40,000 kg/cm² level to which Bridgman's measurements on solid organic compounds were carried. This means that there is a fairly wide pressure range in which these compounds follow the normal compression pattern. The following comparison of theoretical and observed volume ratios for benzene and some of its polynuclear derivatives gives an indication of how the elimination of the excess volume progresses. A measured ratio lower than the theoretical means that some of the excess volume is eliminated in the pressure range for which the ratio is measured, and the amount

of the difference is an indication of the amount by which the normal loss of volume due to compression is increased.

Benzene					
P M kg/cm ²	Ratio			Ratio ⁴⁰ / ₂₅	
	Calc.	Obs.		Calc.	Obs.
⁴⁰ / ₂₀	.938	.920	Benzene	.954	.943
⁴⁰ / ₂₅	.954	.943	Naphtalene	.954	.950
⁴⁰ / ₃₀	.970	.964	Anthracene	.954	.953
⁴⁰ / ₃₅	.985	.984			

As these figures indicate, benzene is just getting rid of the last of the excess volume at the pressure limit of the experiments, and there is no linear section of the benzene compression curve on which the slope can be measured for comparison with the theoretical value. With increased molecular complexity, however, the linear section of the curve lengthens, and for compounds with characteristics similar to those of anthracene there is a 15,000 kg/cm² interval in which the measured volumes should follow the theoretical line.

TABLE 18: *MEASURED VOLUME RATIO - ⁴⁰/₂₅ M kg/cm²*

(Theoretical ratio: .954)

Urea	.954	p-Nitroiodobenzene	.955
Nitrourea	.956	o-Chlorobenzoic acid	.954
Cyanamide	.953	m-Chlorobenzoic acid	.953
o-Xylene	.956	p-Chlorobenzoic acid	.954
p-Xylene	.956	o-Bromobenzoic acid	.954
Triphenyl methane	.953	m-Bromobenzoic acid	.954
o-Diphenyl benzene	.954	p-Bromobenzoic acid	.954
m-Diphenyl benzene	.955	m-Iodobenzoic acid	.955
p-Diphenyl benzene	.955	p-Iodobenzoic acid	.953
Chlorobenzene	.954	p-Nitroaniline	.954
o-Nitrochlorobenzene	.956	o-Acetyl tuluidine	.954
o-Nitrobromobenzene	.955	Tetrahydronaphthalene	.953
p-Nitrobromobenzene	.953	Anthracene	.953
o-Nitroiodobenzene	.953	Acenaphthene	.955

Compounds of this nature have magnetic rotation 3-3 and electric rotation 4. The effective value of s_0^3 is therefore 0.812, and where the compression factors are 4-1½-1 the resulting internal pressure is 127.2 M kg/cm². As shown in the values tabulated for benzene, which were computed on the basis of this internal pressure, the ratio of the volume at 40,000 kg/cm² to that at 25,000 kg/cm² should be 0.954 for all organic compounds with characteristics (molecular complexity, melting point, compression factors, etc.) similar to those of anthracene. Table 18 shows that this theoretical conclusion is corroborated by Bridgman's measurements.

At the time the theoretical values listed in the foregoing tables were originally calculated, Bridgman's results constituted almost the whole of the experimental data then available in the high pressure range, and his experimental limit at 100,000 kg/cm² was the boundary of the empirical knowledge of the effect of high pressure. In the meantime the development of shock wave techniques by American and Russian investigators has enabled measuring compressions at pressures up to several million atmospheres. With the benefit of these new measurements we are now able to extend the correlation between theory and experiment into the region of the maximum compression factors.

The nature of the response of the compression factors to the application of pressure has already been explained, and the maximum factors for each group of elements have been identified. However, the magnitude of the base volume (s_0^3) also enters into the determination of the internal pressure, and coincidentally with the increase in these factors there is a trend toward a minimum base volume. In themselves, modifications of the crystal structure play only a small part in the compressibility picture. Application of sufficient pressure causes a solid to assume one of the crystal forms corresponding to the closest packing of the atoms, the face-centered cubic or close-packed hexagonal for isometric crystals, and the nearest equivalent structures if the crystals are anisometric. If some different crystal form exists at zero pressure, the volume decrease due to the change to one of the close-packed forms shows up as a percentage reduction in all subsequent volumes, but the compressibility is not otherwise affected. However, a difference in crystal structure often indicates a difference in the relative orientation of the atomic rotations. Any such

change in orientation alters the internal pressure, and consequently has a significant effect on the compressibility.

Application of pressure tends to favor what may be called “regular” structures at the expense of those structures that are able to exist only because of special conditions applicable to the particular elements involved. This tendency is evident from the start of the compression process, and is responsible for the large number of deviations from the Chapter 2 values of the inter-atomic distances that are identified by asterisks in Table 14. For example the five elements from chromium to nickel have a number of different inter-atomic distances at low pressure, and are able to crystallize in alternate forms. In the early stages of compression, however, all of these elements, except manganese, orient themselves on the basis of the neutral relative rotation 10, and have an internal pressure that reflects the corresponding value of s_0^3 , which is 0.603. At still higher pressures vanadium shifts to the same relative rotation and joins the group. Manganese probably does likewise, but empirical confirmation of this change is still lacking. Thus the chance of variation of the atomic arrangements is greatly reduced by external pressure. One of the collateral effects is that the amount of uncertainty in the identification of the rotation orientation, and the resulting base volume, is minimized.

Most of the elements that change to a lower base volume at the start of compression maintain this new value of s_0^3 throughout the remainder of the present range of the shock wave experiments. Those that do not make this change in the early stages of compression generally do so at some higher pressure. Only a few keep the same base volume up to the shock wave pressure limit. Still fewer undergo a second transition to a lower base volume. Thus the general pattern involves one reduction of the base volume in the pressure range from zero external pressure up to the limit of the shock wave experiments. This pattern is reflected in the twelve series of measurements that have been selected for comparison with the theoretical values. Out of the twelve elements that are included, only two, copper and chromium, have the same base volume in the shock wave range as at zero pressure. Four continue with the values of s_0^3 applicable to the early stages of compression, the values listed in Table 14, and six change to a lower base volume somewhere above Bridgman’s

pressure limit. The minimum base volumes, the corresponding maximum compression factors, and the resulting internal pressures for these elements are shown in Table 19.

TABLE 19: *MAXIMUM INTERNAL PRESSURES*

	c	a-b	s_0^3	a-z-y	P_0		c	a-b	s_0^3	a-z-y	P_0
V	10	4-3	0.603	4-8-2	1816	Ag	8-10	4-4	0.823	4-8-4	2661
Cr	10	4-3	0.603	4-8-3	2724	W	10	4-4½	0.822	4-8-5	3330
Co	10	4-3	0.603	4-8-3	2724	Au	10	4-4½	0.822	4-8-5	3330
Ni	10	4-3	0.603	4-8-3	2724	Tl	5-10	4-4½	1.074	4-8-5	2549
Cu	8-10	4-3	0.652	4-8-3	2519	Pb	5-10	4-4½	1.074	4-8-5	2549
Mo	10	4-4	0.764	4-8-4	2866	Th	5	4½-4½	1.631	4-8-5	1678

Here again, as in the pressure range of the Bridgman experiments, the theoretical development is not yet far enough advanced to enable specifying the exact locations of the upper sections of the compression curves. Nor is it yet clear in all cases just how many of the possible intermediate values of the compression factors are actually utilized as the pressure increases. What we are able to do at the present rather early stage of the development of the theory is to demonstrate that in this extreme high pressure range, as well as at the lower pressures of the preceding tables, the volume varies inversely with the square root of the total pressure, strictly in accordance with the theory. In this connection it should be noted that the section of each compression curve that is based on the maximum value of the internal pressure is long enough to make the square root pattern clear and distinct.

Furthermore, we are able to show that the slope of the last section of the experimental curve for each element is identical with the theoretical slope determined by the calculated maximum values of the internal pressure, and that the slope of each of the intermediate sections is in agreement with one of the possible intermediate values of that internal pressure. An exact theoretical definition of the curves will have to wait for further progress along the lines discussed earlier. In the meantime, the amount of theoretical information already available will serve as a means of testing the validity of each set of empirical results, and will also enable a reasonable amount of extrapolation of the compression curves beyond the present limits of the shock wave technology.

TABLE 20: *SHOCK WAVE COMPRESSIONS*

P	a-z-y	Calc.	Obs.	a-z-y	Calc.	Obs.	a-z-y	Calc.	Obs.
		W			Au			Mo	
0.1	4-8-3	.972	.970	4-8-1½	.946	.953	4-8-2	.966	.966
.2		.946	.944	4-8-3	.911	.917		.936	.937
.3		.922	.921		.888 R	.888		.908	.912
.4		.900	.901		.867	.864	4-8-3	.885	.890
.5	4-8-4	.880	.882		.847	.843		.868	.870
.6		.865	.866		.828	.825		.851	.852
.7		.850	.851		.811	.810		.966	.836
.8		.836 R	.836		.794	.796		.936	.821
.9		.823	.824	4-8-5	.780	.783		.908	.807
1.0		.810	.812		.711	.772		.795 R	.795
1.1		.798	.800		.762 R	.762		.783	.783
1.2	4-8-5	.787	.790		.754	.752		.771	.772
1.3		.778	.780		.745	.743		.761	.762
1.4		.770	.771		.737	.735		.752	.752
1.5		.762 R	.762		.730	.728		.743 R	.743
1.6		.754	.754		.722	.720		.734	.734
1.7		.747	.746		.715	.714		.726	.726
1.8		.739	.738		.708	.708			
1.9		.732	.731		.701	.702			
2.0		.725	.725		.694	.696			
2.1		.718	.718						
		Cr			Pb			V	
0.1	4-8-1½	.955 R	.955	4-4-1½	.858	.865	4-8-1	.939	.945
.2		.924	.920	4-4-3	.796 R	.796	4-8-1½	.900	.902
.3		.895	.891		.753	.751		.867 R	.867
.4		.869	.867		.716	.718		.838	.838
.5		.845	.846	4-8-3	.691	.693		.811	.812
.6		.823	.827		.673 R	.673		.787	.790
.7	4-8-4	.805	.811		.656	.656		.765	.770
.8		.794	.797		.640	.642	4-8-2	.750	.753
.9		.783	.784	4-8-5	.628	.630		.736	.737
1.0		.772 R	.772		.619 R	.619		.723 R	.723
1.1		.762	.761		.610	.609		.710	.709
1.2	4-8-5	.752	.751		.602	.600		.698	.697
1.3		.742	.742		.594	.593		.687	.686
1.4		.733	.733		.586	.586			

P	a-z-y	Calc.	Obs.	a-z-y	Calc.	Obs.	a-z-y	Calc.	Obs.
		Co			Ni			Cu	
0.1	4-8-1½	.953	.956	4-8-1½	.953	.954	4-8-1	.945	.940
.2		.921	.920		.921	.919		.898	.897
.3		.893	.890		.893	.889	4-8-1½	.865	.864
.4		.867	.865		.867	.865		.838	.836
.5		.843 R	.843		.843 R	.843		.814 R	.814
.6		.821	.823		.821	.825		.792	.794
.7	4-8-4	.801	.806		.801	.808	4-8-3	.772	.777
.8		.782	.791	4-8-3	.790	.794		.760	.762
.9		.769	.776		.779	.780		.749	.749
1.0		.759	.764		.768	.768		.738	.737
1.1		.749	.752		.758	.757		.728	.726
1.2	4-8-5	.739	.741		.748	.747		.718	.716
1.3		.730	.731		.739	.738		.708	.707
1.4		.721	.721		.730	.729		.699	.698
1.5		.712 R	.712		.721 R	.721		.690 R	.690
1.6		.704	.704						
		Ag			Tl			Th	
0.1	4-8-1	.922	.929	4-4-3	.850	.853	4-8-1	.869	.870
.2	4-8-2	.879	.881		.787	.783	4-8-2	.792	.795
.3		.848	.845		.736 R	.736		.474	.744
.4		.820	.817	4-8-3	.702	.703		.710	.707
.5		.749 R	.794		.678 R	.678		.677 R	.677
.6		.771	.775		.656	.658	4-8-3	.652	.652
.7	4-8-4	.752	.759		.637	.642		.632 R	.632
.8		.741	.744	4-8-5	.623	.628		.613	.614
.9		.730	.731		.614	.616		.596	.599
1.0		.749 R	.720		.605 R	.605	4-8-5	.583	.585
1.1		.710	.710		.597	.596		.572	.573
1.2		.701	.700		.588	.587		.562	.562
1.3		.692	.692		.581	.580		.553	.553
1.4		.683	.684		.573	.573		.544	.544
1.5		.675	.677		.566	.567		.535 R	.535
1.6		.667	.670						

Table 20 is a comparison of the theoretical volumes, based on an empirical reference volume for each of the sections of the curves, as in the preceding tables, with the shock wave results obtained at Los Alamos⁵ on the elements that were investigated over the widest range of pressures. Unless there is an increase in the compression factors in the vicinity of 100,000 atmospheres, the compression curves established on the basis of Bridgman's measurements extend into the lower range of these shock wave experiments. In these cases the theoretical volumes up to the first change in the compression factors are calculated on the basis of the reference volume selected from the Bridgman data, and no reference point is identified in this table

A rather surprising feature of these comparisons is that the agreement between the shock wave results and the theoretical volumes is as close as the agreement between Bridgman's static values and the theory. It is true that this set of measurements was deliberately selected for the comparison, and it represents the best results rather than the average, but in any event the close correlation is a significant confirmation of the validity of both the shock wave techniques and the theoretical relations.

The question that now arises is what course the compressibility follows beyond the pressure range of this table. In some cases a transition to a smaller base volume appears to be possible. Copper, for instance, may shift to the rotations of the preceding electropositive elements at some pressure above that of the tabulation. Aside from such special cases, the factors that determine the compressibility in the range below two million atmospheres have reached their limits. At the present stage of the investigation, however, the possibility that some new factor may enter into the picture at extreme pressures cannot be excluded. A "collapse" of the atomic structure of the kind envisioned by the nuclear theory is, of course, impossible, but as matters now stand we are not in a position to say that all aspects of the compressibility situation have been explored. It is conceivable that there may be some, as yet unknown, capability of change in the atomic motions that would increase the resistance to pressure beyond what now appears to be the ultimate limit.

Some shock wave measurements have been made at still higher pressure levels, and these should throw some light on the question. Unfortunately, however, the results are rather ambiguous. Three of the

elements included in these experiments, lead, tin, and bismuth, follow the straight line established in Table 20 up to the maximum pressures of about four million atmospheres. On the other hand, five elements on which measurements were carried to maximums between three and five million atmospheres show substantially lower compressions than a projection of the Table 20 curves would indicate. The divergence in the case of gold, for example, is almost eight percent. But there are equally great differences between the results of different experiments, notably in the case of iron. Whether or not some new factor enters into the compression situation at pressures above those of Table 20 will therefore have to be regarded as an open question.

CHAPTER 5

Heat

If an atom is subjected to an external force of a transient nature, such as that involved in a violent contact, a motion is imparted to it. Where the magnitude of the force is great enough the atom is ejected from the time region and the inter-atomic equilibrium is destroyed. If the force is not sufficient to accomplish this ejection, the motion is turned back at some intermediate point, and it becomes a vibratory, or oscillating, motion. We will identify this oscillation as *thermal motion*.

Where two or more atoms are combined into a molecule, the molecule becomes the thermal unit. The statements about atoms in the preceding paragraph are equally applicable to these molecular units. In order to avoid continual repetition of the expression “atoms and molecules,” the references to thermal units in the discussion that follows will be expressed in terms of molecules, except where we are dealing specifically with substances such as aggregates of metallic elements, in which the thermal units are definitely single atoms. Otherwise the individual atoms will be regarded, for purposes of the discussion, as monatomic molecules.

The thermal motion is something quite different from the familiar vibratory motions of our ordinary experience. In these vibrations that we encounter in everyday life, there is a continuous shift from kinetic to potential energy, and vice versa, which results in a periodic reversal of the direction of motion. In such a motion the point of equilibrium is fixed, and is independent of the amplitude of the vibration. In the thermal situation, however, any motion that is inward in the context of the fixed reference system is coincident with the progression of the natural reference system, and it therefore has no physical effect. Motion in the outward direction is physically effective. From the physical standpoint, therefore, the thermal motion is a net outward motion that

adds to the gravitational motion (which is outward in the time region) and displaces the equilibrium point in the outward direction.

In order to act in the manner described, coinciding with the progression of the natural reference system during the inward phase of the thermal cycle and acting in conjunction with gravitation in the outward phase, the thermal vibration must be a scalar motion. Here again, as in the case of the vibratory motion of the photons, the only available motion form is simple harmonic motion. The thermal oscillation is identical with the oscillation of the photon except that its direction is collinear with the progression of the natural reference system rather than perpendicular to it. However, the suppression of the physical effects of the vibration during the half of the cycle in which the thermal motion is coincident with the reference system progression gives this motion the physical characteristics of an intermittent unidirectional motion, rather than those of an ordinary vibration. Since the motion is outward during half of the total cycle, each natural unit of thermal vibration has a net effective magnitude of one half unit.

Inasmuch as the thermal motion is a property of the individual molecule, not an aspect of a relation between molecules, the factors that come into play at distances less than unity do not apply here, and the direction of the thermal motion, in the context of a stationary reference system is always outward. As indicated earlier, therefore, continued increase in the magnitude of the thermal motion eventually results in destruction of the inter-atomic force equilibrium and ejection of the molecule from the time region. It should be noted, however, that the gravitational motion does not contribute to this result, as it changes direction at the unit boundary. The escape cannot be accomplished until the magnitude of the thermal motion is adequate to achieve this result unassisted.

When a molecule acquires a thermal motion it immediately begins transferring this motion to its surroundings by means of one or more of several processes that will be considered in detail at appropriate points later in this and the subsequent volumes. Coincident with this outflow there is an inflow of thermal motion from the environment, and, in the absence of an externally maintained unbalance, an equilibrium is ultimately reached at a point where inflow and outflow are

equal. Any two molecules or aggregates that have established such an equilibrium with each other are said to be at the same *temperature*.

In the universe of motion defined by the postulates of the Reciprocal System speed and energy have equal standing from the viewpoint of the universe as a whole, but on the low speed side of the neutral axis, where all material phenomena are located, energy is the quantity that exceeds unity. Equality of motion in the material sector is therefore synonymous with equal energy. Thus a temperature equilibrium is a condition in which inflow and outflow of energy are equal. Where the thermal energy of a molecule is fully effective in transfer on contact with other units of matter, its temperature is directly proportional to its total thermal energy content. Under these conditions,

$$E = kT \quad (5-1)$$

In natural units the numerical coefficient k is eliminated, and the equation becomes

$$E = T \quad (5-2)$$

Combining equation 5-2 with equation 4-3 we obtain the *general gas equation*, $PV = T$, or in conventional units,

$$PV = RT \quad (5-3)$$

where R is the *gas constant*.

These are the relations that prevail in the “ideal gas state.” Elsewhere the relation between temperature and energy depends on the characteristics of the transmission process. *Radiation* originates three-dimensionally in the time region, and makes contact one-dimensionally in the outside region. It is thus four-dimensional, while temperature is only one-dimensional. We thus find that the energy of radiation is proportional to the fourth power of the temperature.

$$E_{\text{rad}} = kT^4 \quad (5-4)$$

This relation is confirmed observationally.

The thermal motion originating inside unit distance is likewise four-dimensional in the energy transmission process. However, this motion is not transmitted directly into the outside region in the manner of radiation. The transmission is a contact process, and is subject to the general inter-regional relation previously explained, due to which the value of E' , the energy in the time region, is E^2 as

measured in terms of E . Thus, instead of $E' = k'T^4$, as in radiation, the thermal motion is $E^2 = k'T^4$, or

$$E = kT^2 \quad (5-5)$$

A modification of this relation results from the distribution of the thermal motion over three dimensions of time, while the effective component in thermal interchange is only one-dimensional. This is immaterial as long as the thermal motion is confined to a single rotational unit, but the effective component of the thermal motion of magnetic rotational displacement n is only $1/n^3$ of the total. We may therefore generalize equation 5-5 by applying this factor. Substituting the usual term *heat* (symbol H) for the time region thermal energy E , we then have

$$H = \frac{T^2}{n^3} \quad (5-6)$$

The general treatment of heat in conventional physical theory is empirically based, and is not significantly affected by the new theoretical development. It will not be necessary, therefore, to give this subject matter any attention in this present work, where we are following a policy of not duplicating information that is available elsewhere, except to the extent that reference to such information is required in order to avoid gaps in the theoretical development. The thermal characteristics of individual substances, on the other hand, have not been thoroughly investigated. Since they are of considerable importance, both from the standpoint of practical application and because of the light that they can shed on fundamental physical relationships, it is appropriate to include some discussion of the status of these items in the universe of motion. One of the most distinctive thermal properties of matter is the *specific heat*, the heat increment required to produce a specific increase in temperature. This can be obtained by differentiating equation 5-6.

$$\frac{dH}{dT} = \frac{2T}{n^3} \quad (5-7)$$

Inasmuch as heat is merely one form of energy it has the same natural unit as energy in general, 1.4918×10^{-3} ergs. However, it is more commonly measured in terms of a special heat energy unit, and for present purposes *the natural unit of heat* will be expressed

as 3.5636×10^{-11} gram-calories, the equivalent of the general energy unit.

Strictly speaking, the quantity to which equation 5-7 applies is the specific heat at zero pressure, but the pressures of ordinary experience are very low on a scale where unit pressure is over fifteen million atmospheres, and the question as to whether the equation holds good at all pressures, an issue that has not yet been investigated theoretically, is of no immediate concern. We can take the equation as being applicable under any condition of constant pressure that will be encountered in practice.

The natural unit of specific heat is one natural unit of heat per natural unit of temperature. The magnitude of this unit can be computed in terms of previously established quantities, but the result cannot be expressed in terms of conventional units because the conventional temperature scales are based on the properties of water. The scales in common use for scientific purposes are the *Celsius* or *Centigrade*, which takes the ice point as zero, and the *Kelvin*, which employs the same units but measures from absolute zero. All temperatures stated in this work are absolute temperatures, and they will therefore be stated in terms of the Kelvin scale. For uniformity, the Kelvin notation ($^{\circ}\text{K}$, or simply K) will also be applied to temperature differences instead of the customary Celsius notation ($^{\circ}\text{C}$).

In order to establish the relation of the Kelvin scale to the natural system, it will be necessary to use the actual measured value of some physical quantity involving temperature, just as we have previously used the Rydberg frequency, the speed of light, and Avogadro's number to establish the relations between the natural and conventional units of time, space, and mass. The most convenient empirical quantity for this purpose is the *gas constant*. It will be apparent from the facts developed in the discussion of the gaseous state in a subsequent volume of this series that the gas constant is the equivalent of two-thirds of a natural unit of specific heat. We may therefore take the measured value of this constant, 1.9869 calories, or 8.31696×10^7 ergs, per gram mole per degree Kelvin, as the basis for conversion from conventional to natural units. This quantity is commonly represented by the symbol R , and this symbol will be employed in the conventional manner in the following pages. It should be kept in mind that $R = \frac{2}{3}$ natural unit.

For general purposes the specific heat will be expressed in terms of calories per gram mole per degree Kelvin in order to enable making direct comparisons with empirical data compiled on this basis, but it would be rather awkward to specify these units in every instance, and for convenience only the numerical values will be given. The foregoing units should be understood.

Dividing the gas constant by Avogadro's number, 6.02486×10^{23} per g-mole, we obtain the *Boltzmann constant*, the corresponding value on a single molecule basis: 1.38044×10^{-16} ergs/deg. As indicated earlier, this is two-thirds of the natural unit, and *the natural unit of specific heat* is therefore 2.07066×10^{-16} ergs/deg. We then divide unit energy, 1.49175×10^{-3} ergs, by this unit of specific heat, which gives us 7.20423×10^{12} degrees Kelvin, *the natural unit of temperature* in the region outside unit distance (that is, for the gaseous state of matter).

We will also be interested in the unit temperature on the T^3 basis, the temperature at which the thermal motion reaches the time region boundary. The $\frac{3}{4}$ power of 7.20423×10^{12} is 4.39735×10^9 . But the thermal motion is a motion of matter and involves the $\frac{2}{9}$ vibrational addition to the rotationally distributed linear motion of the atoms. This reduces the effective temperature unit by the factor $1 + \frac{2}{9}$, the result being 3.5978×10^9 degrees K.

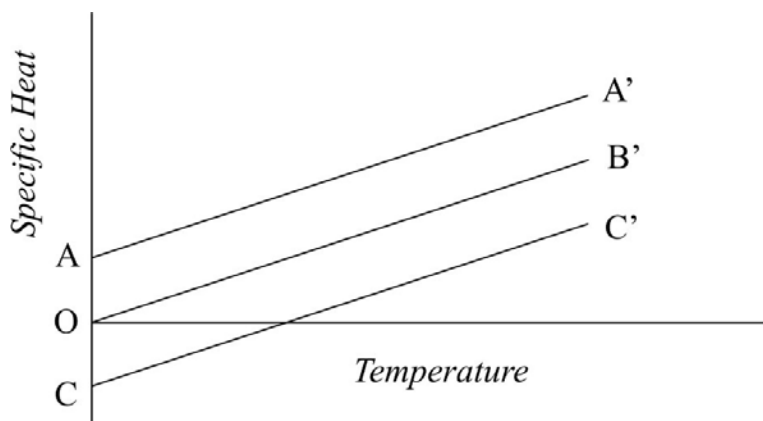
On first consideration, this temperature unit may seem incredibly large, as it is far above any observable temperature, and also much in excess of current estimates of the temperatures in the interiors of the stars, which, according to our theoretical findings, can be expected to approach the temperature unit. However, an indication of its validity can be obtained by comparison with the unit of pressure, inasmuch as the temperature and pressure are both relatively simple physical quantities with similar, but opposite, effects on most physical properties, and should therefore have units of comparable magnitude. The conventional units, the degree K and the gram per cubic centimeter have been derived from measurements of the properties of water, and are therefore approximately the same size. Thus the ratio of natural to conventional units should be nearly the same in temperature as in pressure. The value of the temperature unit just calculated, 3.5978×10^9 degrees K, conforms to this theoretical requirement, as the natural unit of pressure derived in Volume I is 5.386×10^9 g/cm³.

Except insofar as it enters into the determination of the value of the gas constant, the natural unit of temperature defined for the gaseous state plays no significant role in terrestrial phenomena. Here the unit with which we are primarily concerned is that applicable to the condensed states. Just as the gaseous unit is related to the maximum temperature of the gaseous state, the lower unit is related to the maximum temperature of the liquid state. This is the temperature level at which the unit molecule escapes from the time region in one dimension of *space*. The motion in this low energy range takes place in only one scalar dimension. We therefore reduce the three-dimensional unit, 3.5978×10^9 K, to the one-dimensional basis, and divide it by 3 because of the restriction to one dimension of space. The natural unit applicable to the condensed state is then $\frac{1}{3}(3.598 \times 10^9)^{1/2}$ R degrees K = 510.7 K.

The magnitude of this unit was evaluated empirically in the course of a study of liquid volume carried out prior to the publication of *The Structure of the Physical Universe* in 1959. The value derived at that time was 510.2, and this value was used in a series of articles on the liquid state that described the calculation of the numerical values of various liquid properties, including volume, viscosity, surface tension, and the critical constants. Both the 510.2 liquid unit and the gaseous unit were listed in the 1959 publication, but the value of the gaseous unit given there has subsequently increased by a factor of 2 as a result of a review of the original derivation.

Since the basic linear vibrations (photons) of the atom are rotated through all dimensions, they have active components in the dimensions of any thermal motion, whatever that dimension may be, just as they have similar components parallel to the rotationally distributed motions. As we found in our examination of the effect on the rotational situation, this basic vibrational component amounts to $\frac{2}{9}$ of the primary magnitude. Because the thermal motion is in time (equivalent space) its scalar direction is not fixed relative to that of the vibrational component. This vibrational component will therefore either supplement or oppose the thermal specific heat. The net specific heat, the measured value, is the algebraic sum of the two. The presence of the vibrational component does not change the linear relation of the specific heat to the temperature, but it does alter the zero point, as indicated in Fig. 2.

FIGURE 2



In this diagram the line OB' is the specific heat curve derived from equation 5-7, assuming a constant value of n and a zero initial level. If the scalar direction of the vibrational component is opposite to that of the thermal motion, the initial level is positive; that is, a certain amount of heat must be supplied to neutralize the vibrational energy before there is *any* rise in temperature. In this case the specific heat follows the line AA' parallel to OB' and above it. If the scalar direction of the vibrational component is the same as that of the thermal motion, the initial level is negative, and the specific heat follows the line CC' , likewise parallel to OB' but below it. Here there is an effective temperature due to the vibrational energy *before* any thermal motion takes place.

Although this initial component of the molecular motion is *effective* in determining the temperature, its magnitude cannot be altered and it is therefore not *transferable*. Consequently, even where the initial level is negative, there is no negative specific heat. Where the sum of the negative initial level and the thermal component is negative, the effective specific heat of the molecule is zero.

It should be noted in passing that the existence of this second, fixed, component of the specific heat confirms the vibrational character of the basic constituent of the atomic structure, the constituent that we have identified as a photon. The demonstration that there is a negative initial level of the specific heat curve is a clear indication of the validity of the theoretical identification of the basic unit in the atomic structure as a vibratory motion.

Equation 5-7 can now be further generalized to include the specific heat contribution of the basic vibration: the initial level, which we will represent by the symbol I . The net specific heat, the value as measured, is then

$$\frac{dH}{dT} = \frac{2T}{n^3} + I \quad (5-8)$$

Where there is a choice between two possible states, as there is between the positive and negative initial levels, the probability relations determine which of the alternatives will prevail. Other things being equal, the condition of least net energy is the most probable, and since the negative initial level requires less net energy for a given temperature than the positive initial level, the thermal motion is based on the negative level at low temperatures unless motion on this basis is inhibited by structural factors.

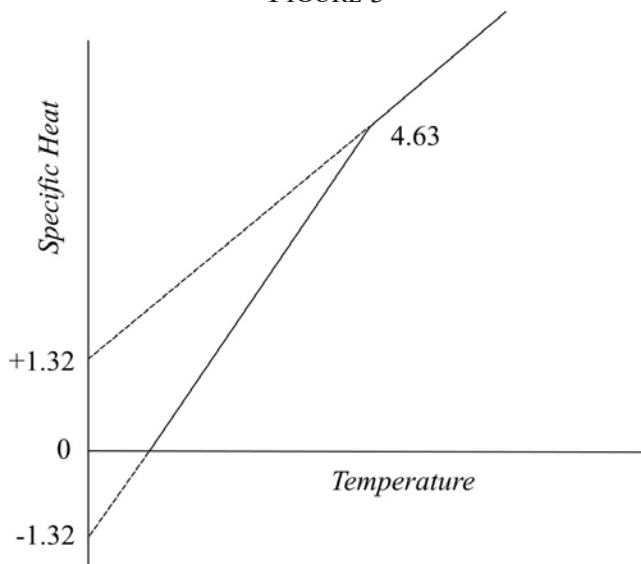
Addition of energy in the time region takes place by means of a decrease in the effective time magnitude, and it involves eliminating successive time units from the vibration period. The process is therefore discontinuous, but the number of effective time units under ordinary conditions is so large that the relative effect of the elimination of one unit is extremely small. Furthermore, observations of heat phenomena of the solid state do not deal with single molecules but with aggregates of many molecules, and the measurements are averages. For all practical purposes, therefore, we may consider that the specific heat of a solid increases in continuous relation to the temperature, following the pattern defined by equation 5-8.

As pointed out earlier in this chapter, the thermal motion cannot cross the time region boundary until its magnitude is sufficient to overcome the progression of the natural reference system without assistance from the gravitational motion; that is, it must attain unit magnitude. The maximum thermal specific heat, the total increment above the initial level, is the value that prevails at the point where the thermal motion reaches this unit level. We can evaluate it by giving each of the terms T and n in equation 5-7 unit value, and on this basis we find that it amounts to 2 natural units, or $3R$. The normal initial level is $-\frac{2}{9}$ and of this $3R$ specific heat, or $-\frac{2}{3}R$. The $3R$ total is then reached at a net positive specific heat of $2\frac{1}{3}R$.

Beyond this $3R$ thermal specific heat level, which corresponds to the regional boundary, the thermal motion leaves the time region and undergoes a change which requires a substantial input of thermal energy to maintain the same temperature, as will be explained later. The condition of minimum energy, the most probable condition, is maintained by avoiding this regional change by whatever means are available. One such expedient, the only one available to molecules in which only one rotational unit is oscillating thermally, is to change from a negative to a positive initial level. Where the initial level is $+\frac{2}{3}R$ instead of $-\frac{2}{3}R$, the net positive specific heat is $3\frac{2}{3}R$ at the point where the thermal specific heat reaches the $3R$ limit. The regional transmission is not required until this higher level is reached. The resulting specific heat curve is shown in Fig. 3.

Inasmuch as the magnetic rotation is the basic rotation of the atom, the maximum number of units that can vibrate thermally is ordinarily determined by the magnetic displacement. Low melting points and certain structural factors impose some further restrictions, and there

FIGURE 3



are a few elements, and a large number of compounds that are confined to the specific heat pattern of Fig. 3, or some portion of it. Where the thermal motion extends to the second magnetic rotational

unit, to *rotation two*, we may say, using the same terminology that was employed in the inter-atomic distance discussion, the Fig. 3 pattern is followed up to the $2\frac{1}{3}$ level. At that point the second rotational unit is activated. The initial specific heat level for rotation two is subject to the same n^3 factor as the thermal specific heat, and it is therefore $\frac{1}{n^3} \times 2\frac{1}{3} R = \frac{1}{12} R$. This change in the negative initial level raises the net positive specific heat corresponding to the thermal value $3R$ from $2.333 R$ to $2.917 R$, and enables the thermal motion to continue on the basis of the preferred negative initial level up to a considerably higher temperature.

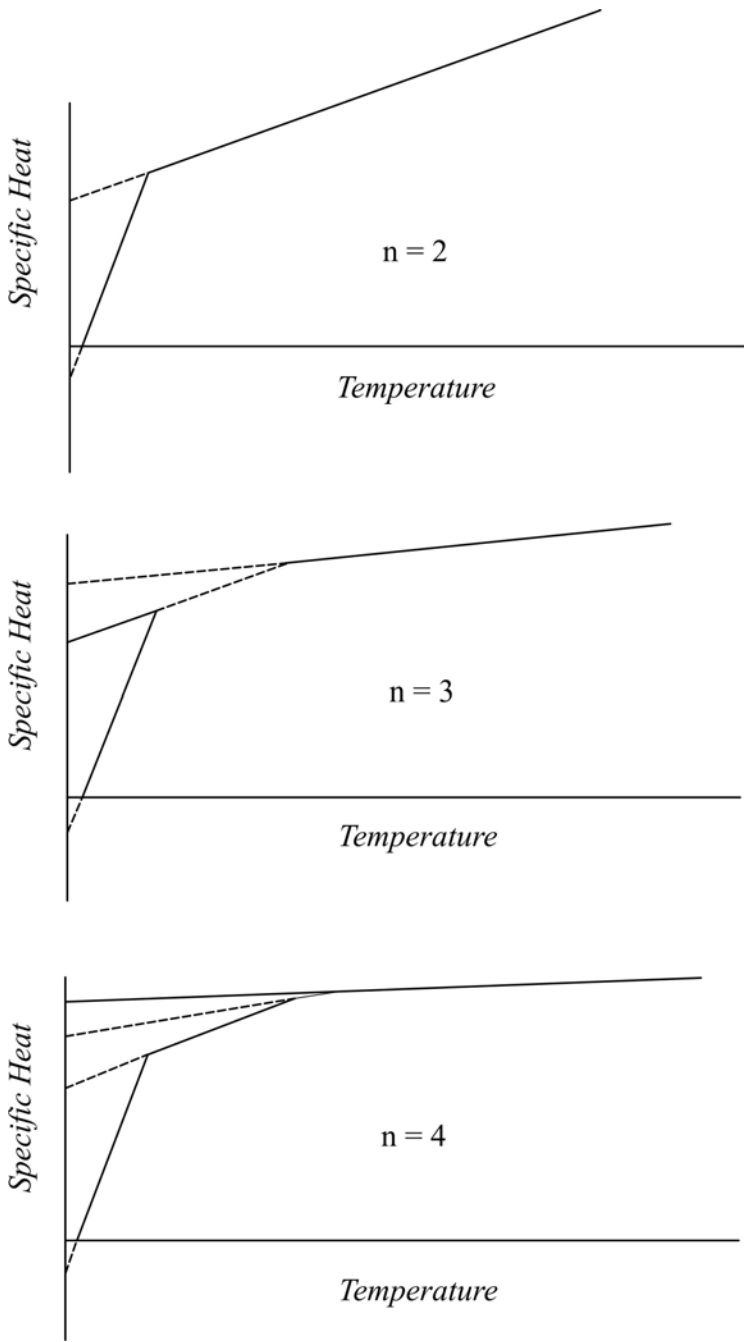
When the rotation two curve reaches its end point at $2.917 R$ net positive specific heat, a further reduction of the initial level by a transition to the rotation three basis, where the higher rotation is available, raises the maximum to $2.975 R$.

Another similar transition follows, if a fourth vibrating unit is possible. The following tabulation shows the specific heat values corresponding to the initial and final levels of each curve. As indicated earlier, the units applicable to the second column under each heading are calories per gram mole per degree Kelvin.

Vibrating Units	Effective Initial Level		Maximum Net Specific Heat (negative initial level)	
1	-0.6667 R	-1.3243	2.3333 R	4.6345
2	-0.0833 R	-0.1655	2.9167 R	5.7940
3	-0.0247 R	-0.0490	2.9753 R	5.9104
4	-0.0104 R	-0.0207	2.9896 R	5.9388

Ultimately the maximum net positive specific heat that is possible on the basis of a negative initial level is attained. Here a transition to a positive initial level takes place, and the curve continues on to the overall maximum. As a result of this mechanism of successive transitions, each number of vibrating units has its own characteristic specific heat curve. The curve for rotation one has already been presented in Fig. 3. For convenient reference we will call this a type two curve. The different type one curves, those of two, three, and four vibrating units, are shown in Fig. 4. As can be seen from these diagrams, there is a gradual flattening and an increase in the ratio of temperature to specific heat as the number of vibratory units increases.

FIGURE 4



The actual temperature scale of the curve applicable to any particular element or compound depends on the thermal characteristics of the substance, but the relative temperature scale is determined by the factors already considered, and the curves in Fig. 4 have been drawn on this relative basis.

As indicated by equation 5-8, the slope of the rotation two segment of the specific heat curve is only one-eighth of the slope of the rotation one segment. While this second segment starts at a temperature corresponding to $2\frac{1}{3}R$ specific heat, rather than from zero temperature, the fixed relation between the two slopes means that a projection of the two-unit curve back to zero temperature always intersects the zero temperature ordinate at the same point regardless of the actual temperature scale of the curve. The slopes of the three-unit and four-unit curves are likewise specifically related to those of the earlier curves, and each of these higher curves also has a fixed initial point. We will find this feature very convenient in analyzing complex specific heat curves, as each experimental curve can be broken down into a succession of straight lines intersecting the zero ordinate at these fixed points, the numerical values of which are as follows:

Vibrating Units	Specific Heat at 0° K (projected)	
1	-0.6667 R	-1.3243
2	1.9583 R	3.8902
3	2.6327 R	5.2298
4	2.8308 R	5.6234

These values and the maximum net specific heats previously calculated for the successive curves enable us to determine the relative temperatures of the various transition points. In the rotation three curve, for example, the temperatures of the first and second transition points are proportional to the differences between their respective specific heats and the 3.8902 initial level of the rotation two segment of the curve, as both of these points lie on this line. The relative temperatures of any other pair of points located on the same straight line section of any of the curves can be determined in a similar manner. By this means the following relative temperatures have been calculated, based on the temperature of the first transition point as unity.

Vibrating Units	Relative Temperature Transition Point	End Point
1	1.000	1.80
2	2.558	4.56
3	3.086	9.32
4	3.391	17.87

The curves of Figs. 3 and 4 portray what may be called the “regular” specific heat patterns of the elements. These are subject to modifications in certain cases. For instance, all of the electronegative elements with displacements below 7 thus far studied substitute an initial level of -0.66 for the normal -1.32 . Another common deviation from the regular pattern involves a change in the temperature scale of the curve at one of the transition points, usually the first. For reasons that will be developed later, the change is normally downward. Inasmuch as the initial level of each segment of the curve remains the same, the change in the temperature scale results in an increase in the slope of the higher curve segment. The actual intersection of the two curve segments involved then takes place at a level above the normal transition point.

There are some deviations of a different nature in the upper portions of the curves where the temperatures are approaching the melting points. These will not be given any consideration at this time because they are connected with the transition to the liquid state and can be more conveniently examined in connection with the discussion of liquid properties.

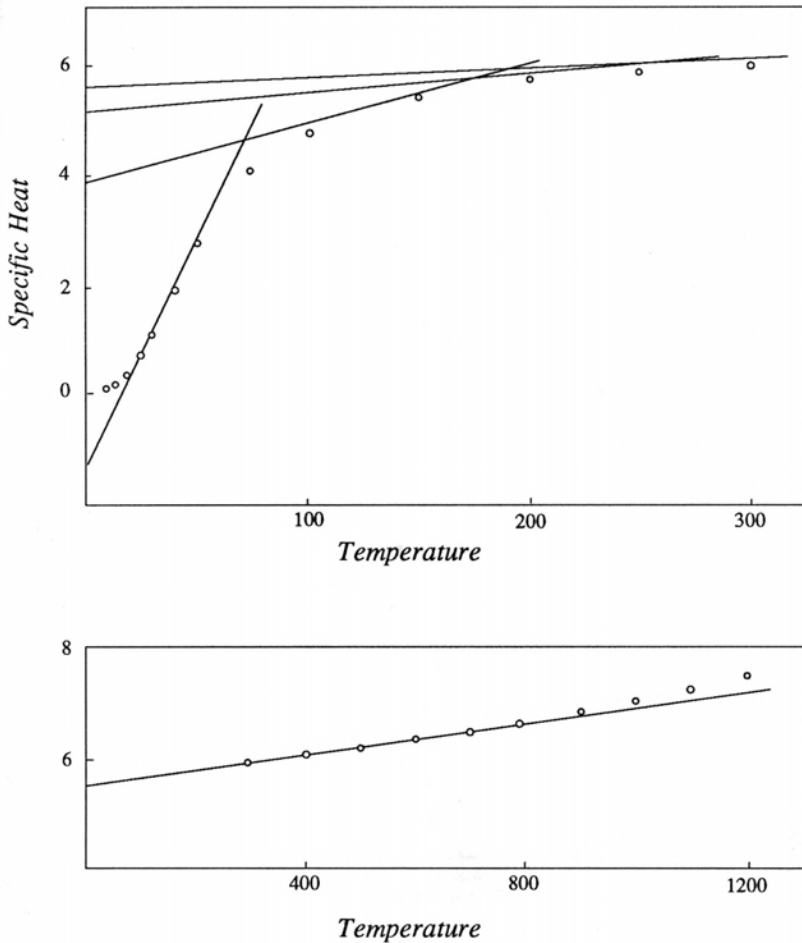
As mentioned earlier, the quantity with which this and the next two chapters are primarily concerned is the specific heat at zero external pressure. In Chapter 6 the calculated values of this quantity will be compared with measured values of the specific heat at constant pressure, as the difference between the specific heat at zero pressure and that at the pressures of observation is negligible. Most conventional theory deals with the specific heat at constant volume rather than at constant pressure, but our analysis indicates that the measurement under constant pressure corresponds to the fundamental quantity.

CHAPTER 6

Specific Heat Patterns

Fig. 5 is a specific heat curve derived from experimental data.

FIGURE 5: *SPECIFIC HEAT – SILVER*



The points shown in this graph are the measured values of the specific heat of silver. The accompanying solid lines are the segments of the theoretical four-unit curve of Fig. 4 with the temperature scale located empirically. While the curve defined by the plotted points has the same general shape as the theoretical curve, it is quite different in appearance inasmuch as the sharp angles of the theoretical curve have been replaced by smooth and gradual transitions.

The explanation of this difference lies in the manner in which the measurements are made. As indicated by equation 5-8 and the curves in Figs. 3 and 4, the specific heat of an individual molecule can be represented by a succession of straight lines. Experimental observations, however, are not made on single molecules, but on aggregates of molecules, and the observed temperature of the aggregate is the average of many different individual molecular temperatures, which are distributed about the average in accordance with the probability relations. Midway between the transition points the relation between temperature and specific heat for most of the individual molecules is such that their specific heats lie on the same straight line in the diagram. The average consequently lies on the same line, and coincides with the true molecular specific heat corresponding to the average temperature. In the neighborhood of a transition point, however, the molecules that are individually at the higher temperatures cannot continue on the same line beyond the $3R$ limit, and must conform to a lower curve based on a higher number of rotating units. This operates to reduce the specific heat of the aggregate below the true molecular value for the prevailing average temperature.

In the silver curve, Fig. 5, for example, the true atomic specific heat at 75 K is 4.69. This would also be the average specific heat of the silver aggregate at that temperature if the silver atoms were able to continue vibrating on the basis of one rotating unit up to the point beyond which the probability distribution is negligible. But at a specific heat of $2\frac{1}{3}R$ (4.633) the vibration changes to the two-unit basis. Those atoms in the probability distribution that have specific heats above this level cannot conform to the one-unit line but must follow a line that rises at a much slower rate. The lower specific heat of these atoms reduces the average specific heat of the aggregate, and causes the aggregate curve to diverge more and more

from the straight line relation as the proportion of atoms reaching the transition point increases. The divergence reaches a maximum at the transition temperature, after which the specific heat of the aggregate gradually approaches the upper atomic curve. Because of this divergence of the measured (aggregate) specific heats from the values applicable to the individual atoms the specific heat of silver at 75 K is 4.10 instead of 4.69.

A similar effect in the opposite direction can be seen at the lower end of the silver curve. Here the specific heat of the aggregate (the average of the individual values) could stay on the one-unit theoretical curve only if it were possible for the individual specific heats to fall below zero. But there is no negative thermal energy, and the atoms which are individually at temperatures below the point where the curve intersects the zero specific heat level all have zero thermal energy and zero specific heat. Thus there is no negative deviation from the average, and the positive deviation due to the presence of atoms with individual temperatures above zero constitutes the specific heat of the aggregate. The specific heat of a silver *atom* at 15 K is zero, but the measured specific heat of a silver *aggregate* at an average temperature of 15 K is 0.163.

Evaluation of the deviations from the linear relationship in these transitional regions involves the application of probability mathematics, the validity of which was assumed as a part of the Second Fundamental Postulate of the Reciprocal System. For reasons previously explained, a full treatment of the probability aspects of the phenomena now under discussion is beyond the scope of this work, but a general consideration of the situation will enable us to arrive at some qualitative conclusions which will be adequate for present purposes.

At the present stage of development of probability theory there are a number of probability functions in general use, each of which seems to have advantages for certain applications. For the purpose of this work the appropriate function is one that expresses the results of pure chance without modification by any other factor. Such a function is strictly applicable only where the units involved are all exactly alike, the distribution is perfectly random, the units are infinitely small, the

variability is continuous, and the size of the group is infinitely large. The ordinary classes of events around which most present-day probability theory has been constructed, such as coin and dice experiments, obviously fail to meet these requirements by a wide margin. Coins, for instance, are not continuously variable with an infinite number of possible states. They have only two states, heads and tails. This means that a major item of uncertainty has become almost a certainty, and the shape of the probability distribution curve has been altered accordingly. Strictly speaking, it is no longer a true probability curve, but a combination curve of probability and knowledge.

The basic physical phenomena do conform closely to the requirements of a system in which the laws of pure chance are valid. The units are nearly uniform, the distribution is random, the variability is continuous, or nearly continuous, and the size of the group, although not infinite, is extremely large. If any of the probability functions in general use can qualify as representing pure chance the most likely prospect is the so-called “normal” probability function, which can be expressed as

$$\frac{1}{y} = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} \quad (6-1)$$

Tables of this function and its integral to fifteen decimal places are available.⁶ It has been found in the course of this work that sufficient accuracy for present purposes can be attained by calculating probabilities on the basis of this expression, and it has therefore been utilized in all of the probability applications discussed herein, without necessarily assuming the absolute accuracy of this function in these applications, or denying the existence of more accurate alternatives. For example, Maxwell’s asymmetric probability distribution is presumably accurate in the applications for which it was devised (a point that has not yet been examined in the context of the Reciprocal System), and it may also apply to some of the phenomena discussed in this work. However, the results thus far obtained, particularly in application to the liquid properties, favor the normal function. In any event it is clear that if any error is introduced by utilizing the normal function it is not large enough to be significant in this first general treatment of the subject matter.

On the foregoing basis, the distribution of molecules with different individual temperatures takes the form of a probability function Φ_t , where t is the deviation from the average temperature. The contribution of the Φ_t molecules at any specified temperature to the deviation of the specific heat from the theoretical value corresponding to the average temperature depends not only on the number of these molecules but also on the magnitude of the specific heat deviation attributable to each molecule; that is, the difference between the specific heat of the molecule and that of a molecule at the average temperature of the aggregate. Since the specific heat segment from which the deviation takes place is linear, this deviation is proportional to the temperature difference t , and may be represented as kt . The total deviation due to the Φ_t molecules at temperature t is then $kt\Phi_t$, and the sum of all deviations in one direction (positive or negative) may be obtained by integration.

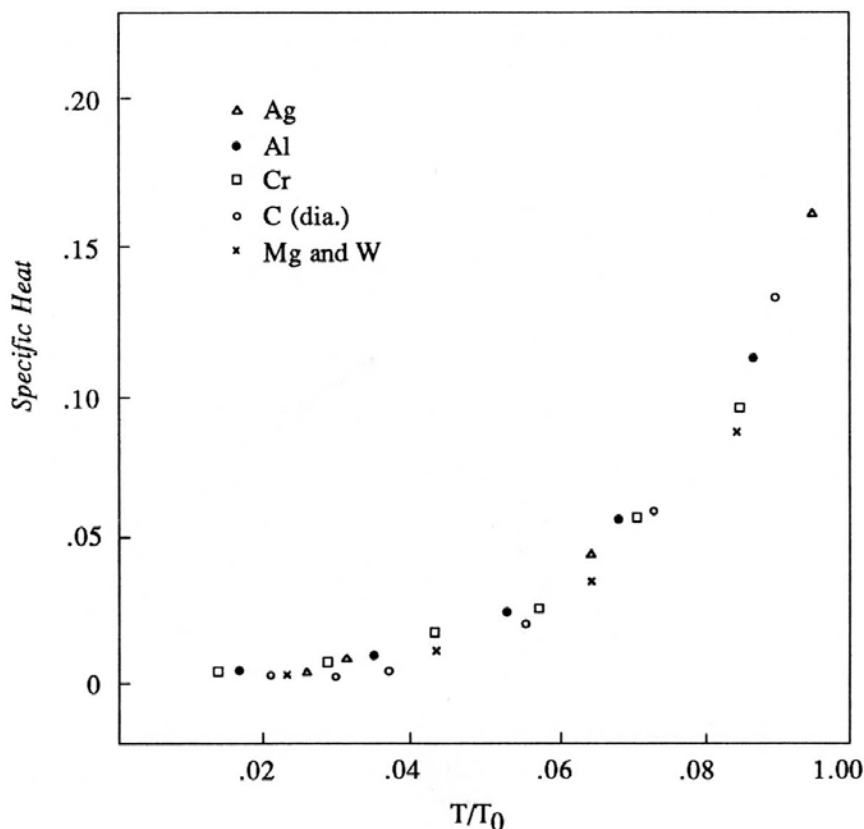
It is quite evident that the deviations of the experimental specific heat curves from the theoretical straight lines, both at the zero level and at the transition point have the general characteristics of the probability curves. However, the experimental values are not accurate enough, particularly in the temperature range of the lower transition, to make it worth while to attempt any quantitative correlations between the theoretical and experimental results. Furthermore, there is still some theoretical uncertainty with respect to the proper application of the probability function that prevents specifying the exact location of the probability curve.

The uncertain element in the situation is the magnitude of the probability unit. Equation 6-1 is complete mathematically, but in order to apply it, or any of its derivatives, to any physical situation it is necessary to ascertain the physical unit corresponding to the mathematical unit. One pertinent question still lacking a definite answer is whether this probability unit is the same for all substances. If so, the lower portion of the curve, when reduced to a common temperature base, should be the same for all substances with the -1.32 initial level. On this basis, the specific heat of the aggregate at the temperature T_0 , where the theoretical curve intersects the zero axis, should be a constant. Actually, most of the elements with the -1.32 initial level do have a measured specific heat in the

neighborhood of 0.20 at this point, but a few others show substantial deviations from this value. It is not yet clear whether this is a result of variability in the probability unit, or reflects inaccuracies in the experimental values.

Whether all of the curves with the same maximum deviation (0.20) are coincident below T_0 is likewise still somewhat uncertain. There is a greater spread in the observed specific heats below 0.20 than can be ascribed to errors in measurement, but most of the scatter can probably be explained as the result of lack of thermal equilibrium. At these low temperatures it no doubt takes a long time to establish equilibrium, and even an accurate measurement will not produce the correct result unless the aggregate is in thermal equilibrium. It is significant that the specific heats of the common elements which

FIGURE 6: *SPECIFIC HEAT - LOW TEMPERATURES*



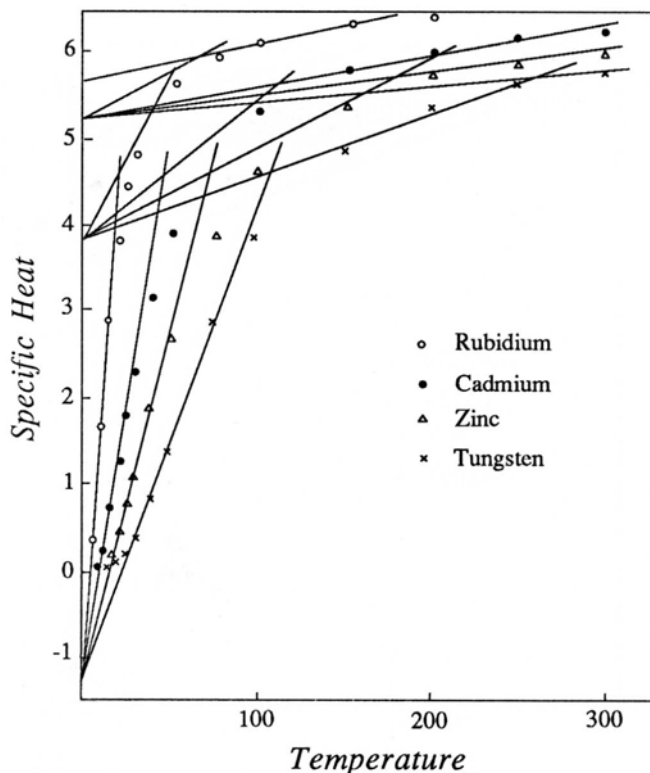
have been studied most extensively deviate only slightly from a smooth curve in this low temperature region. Fig. 6, which shows the measured values of the specific heats of six of these elements on a temperature scale relative to T_0 , demonstrates this coincidence.

If the probability unit is the same for all, or most, of the elements, as these data suggest, the deviation of the experimental curve from the theoretical curve for the single atom at the first transition point, T_1 , should also have a constant value. Preliminary examination of the curves of the elements that follow the regular pattern indicates that the values of this deviation actually do lie within a range extending from about 0.55 to about 0.70. Considerable additional work will be required before these curves can be defined accurately enough to determine whether there is complete coincidence, but present indications are that the deviation at T_1 is, in fact, a constant for all of the regular elements, and is in the neighborhood of three times the deviation at T_0 .

With the benefit of the foregoing information as to the general nature of the deviations from the theoretical curves of Chapter 5 due to the manner in which the measurements are made, we are now prepared to examine the correlation between the theoretical curves and the measured specific heats. In order to arrive at a complete definition of the specific heat of a substance it is not only necessary to establish the shapes of the specific heat curves, the objective at which most of the foregoing discussion is aimed, but also to define the temperature scale of each curve. Although the theoretical conclusions with respect to these two theoretical aspects of the specific heat situation, like all other conclusions in this work, are derived by developing the consequences of the fundamental postulates of the Reciprocal System of theory, they are necessarily reached by two different lines of theoretical development. For this reason a more meaningful comparison with the experimental data can be presented if we deal with these two aspects independently. In this chapter, therefore, the experimental values will be compared graphically with theoretical curves in which the temperature scales are empirical. Chapter 7 will complete the definition of the curves by deriving the relevant temperature magnitudes.

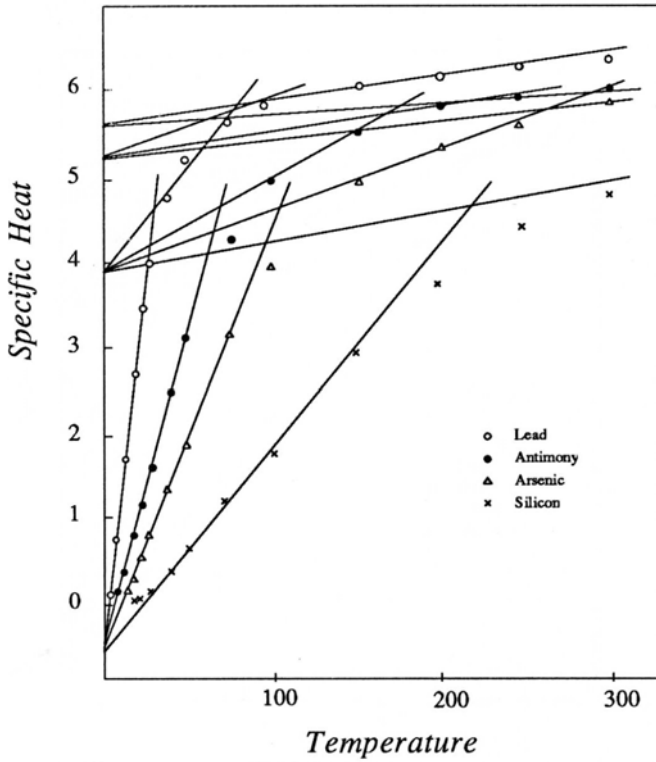
The curves of Fig. 7 are typical of those of most of the elements.⁷ As indicated in Fig. 4, the final straight line segment of each curve

occupies the greater part of the temperature range of the solid state in the case of the high melting point elements. The significant features of the curves are therefore confined to the lower temperatures, and in order to bring them out more clearly only the lower temperature range (up to 300° K) is shown in the illustrations that follow. The

FIGURE 7: *SPECIFIC HEAT*

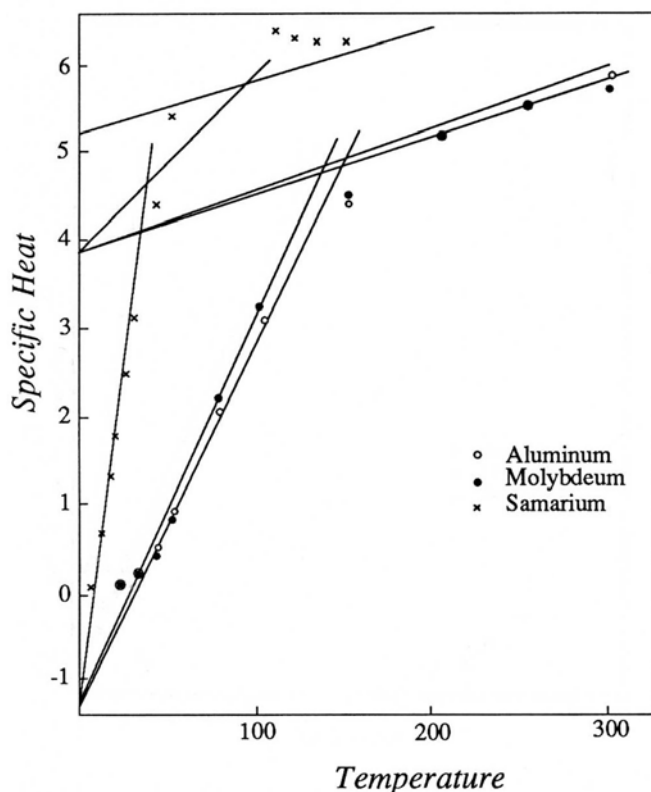
remaining sections of the curves of Fig. 7 are extensions of the lines shown in the diagram, except in the case of tungsten, which undergoes a transition to the four-unit status at about 325 °K.

Fig. 8 is a similar group of specific heat curves for four of the electronegative elements with the -0.66 initial level. Aside from this higher initial level these curves are identical with those of Fig. 7 when all are reduced to a common temperature scale. The transition to the two-unit vibration takes place at 4.63 ($2\frac{1}{3}R$) regardless of the higher initial level. This point will be given further consideration in

FIGURE 8: *SPECIFIC HEAT*

Chapter 7. The upper portions of the lead and antimony curves, which are not shown on the graph, are extensions of the lines in the diagram. Arsenic and silicon have transitions at temperatures above 300°K .

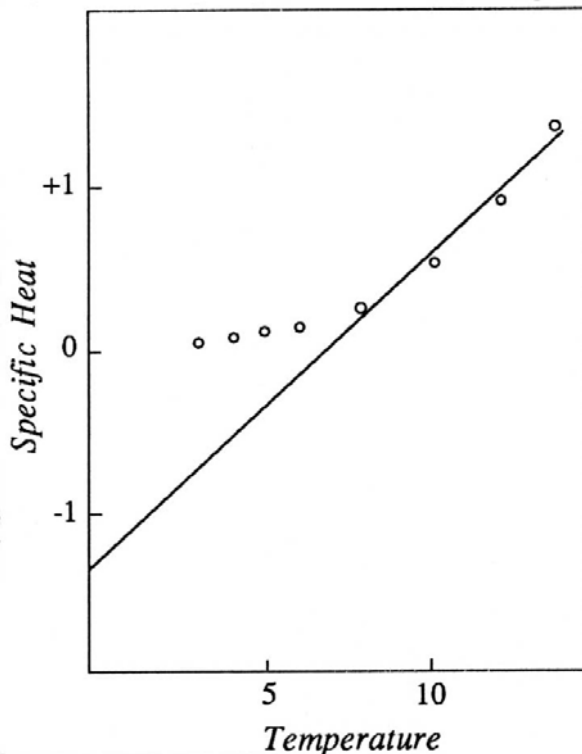
As noted in Chapter 5, there are a number of elements that undergo a modification of the temperature scale at the first transition point. Two curves with the modified second segment are shown in Fig. 9. These two curves actually apply to four elements, as the specific heat of lithium follows the aluminum curve, while that of ruthenium coincides with the molybdenum curve. Coincidence of the specific heat curves of different elements, as in the instances mentioned, is not as uncommon as might be expected. The number of possible curve patterns is quite limited, and, as we will see in the next chapter, where the nature of the change in the temperature will be examined, the temperature factors are confined to specific values mainly within a relatively narrow range.

FIGURE 9: *SPECIFIC HEAT*

Also included in Fig. 9 is an example of a specific heat curve for an element which undergoes an internal rearrangement that modifies the thermal pattern. The measurements shown for samarium follow the regular pattern up to the vicinity of the first transition point at 35 K. Some kind of a modification of the molecular structure is evidently initiated at this point in lieu of, or in addition to, the normal transition to the two-unit vibrational status. This absorbs a considerable quantity of heat, which manifests itself as an addition to the measured specific heat over the next portion of the temperature range. By about 175 K the readjustment is complete, and the specific heat returns to the normal curve. Most of the other rare earth elements undergo similar readjustments at comparable temperatures. Elsewhere, if changes of this kind take place at all, they almost always occur at relatively high temperatures. The reason for this peculiarity of the rare earth group is, as yet, unknown.

All of the types of deviations from the regular pattern that have been discussed thus far are found in the electronegative elements of the lower rotational groups. There is also an additional source of variability in the specific heats of these elements, as their atoms can combine with each other to form molecules. The result is a wide enough variety of behavior to give almost every one of these elements a unique specific heat curve. Of special interest are those cases in which the variation is accomplished by omitting features of the regular pattern. The neon curve, for example, is a single straight line from the -1.32 initial level to the melting point. The specific heat curve of a hydrogen molecule, Fig. 10, is likewise a single straight line, but hydrogen has no rotational specific heat component at all, and this line therefore extends only from the negative initial level, -1.32, to the specific heat of the positive initial level, +1.32, at which point melting takes place.

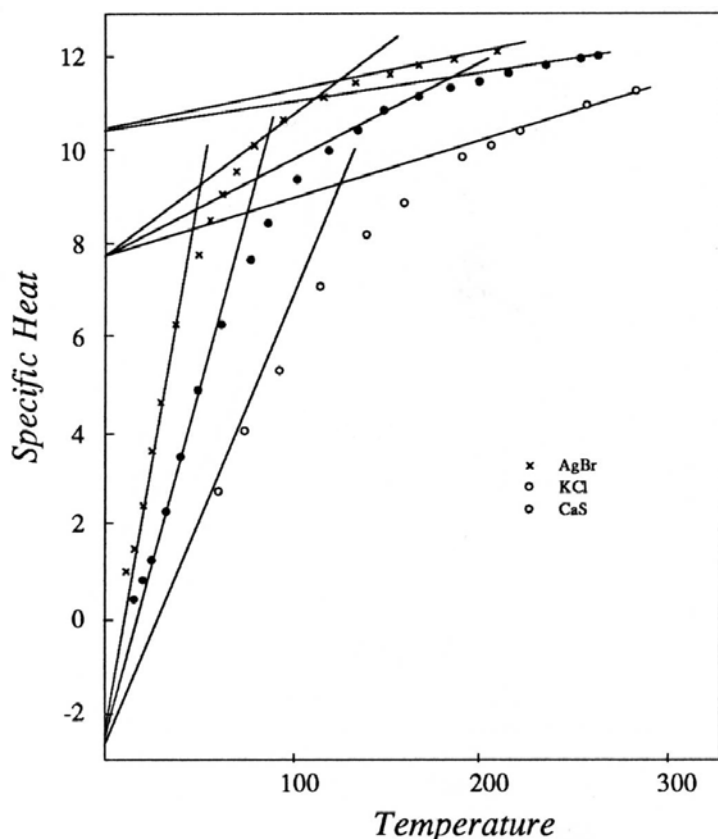
FIGURE 10: *SPECIFIC HEAT – HYDROGEN*



The specific heats of binary compounds based on the normal orientation, simple combinations of Division I and Division IV elements, follow the same pattern as those of the electropositive elements. In these compounds each atom behaves as an individual thermal unit just as it would in a homogeneous aggregate of like atoms. The molecular specific heats of such compounds are twice as large as the values previously determined for the elements, not because the specific heat per atom is any different, but because there are two atoms in each formula molecule.

The curves for KCl and CaS, Fig. 11, illustrate the specific heat pattern of this class of compounds. Some binary compounds of other structural types conform to the same regular pattern as in the curve for AgBr, also shown in Fig.11. As in the elements there is

FIGURE 11



also a variation of this regular pattern in which certain compounds of the electronegative elements have a higher initial level, but in the compounds such as ZnO and SnO this level is zero, rather than -0.66, as it is in the elements.

Some of the larger molecules similarly act thermally as associations of independent atoms. CaF_2 and FeS_2 are typical. More often, however, two or more of the constituent atoms of the molecule act as a single thermal unit. For example, both the KHF_2 molecule, which contains four atoms, and the CsClO_4 molecule, which contains six, act thermally as three units. In the subsequent discussion the term *thermal group* will be used to designate any combination of atoms that acts as a single thermal unit. Where individual atoms participate in thermal motion jointly with groups of atoms, the individual atoms will be regarded as monatomic groups. On this basis we may say that there are three thermal groups in each of the KHF_2 and CsClO_4 molecules.

The great majority of compounds not only form thermal groups but also alter the number of groups in the molecule as the temperature varies. A common pattern is illustrated by the chromium chlorides. CrCl_2 acts as a single thermal group at very low temperatures; CrCl_3 as two. The initial specific heat levels are -1.32 and -2.64 respectively. There is a gradual increase in the average number of thermal groups per molecule up to the first transition point, at which temperature all atoms are acting independently. At the initial point of the second segment of the curve this independent status is maintained, and above the transition temperature the CrCl_2 molecule acts as three thermal groups, while CrCl_3 has four.

At the present stage of the investigation we can determine from theory the possible ways in which a molecule can split up into thermal groups, but we are not yet able to specify on theoretical grounds just which of these possibilities will prevail at any given temperature, or where the transition from one to the other will take place. The theoretical information thus far developed does, however, enable us to analyze the empirical data and to establish the specific heat pattern of each substance; that is, to determine just how it acts thermally. Aside from some cases, mainly involving very large molecules,

where the specific heat pattern is unusually complex, and in those instances where experimental errors lead to erroneous interpretation, it is possible to identify the effective number of thermal groups at the critical points of the curves. Once this information is available for any substance, the definition of its specific heat curve is essentially complete, except for the temperature scale, the determinants of which will be identified in Chapter 7. Where n is the number of active thermal groups in a compound, the initial level is $-1.32 n$, the initial point of the second segment of a Type 1 curve is $3.89 n$, and the first transition point is $4.63 n$.

The tendency of the atoms of multi-atom molecules to form thermal groups is particularly evident where the molecules contain radicals, because of the major differences in the cohesive forces that are responsible for the existence of the radicals. The extent to which the association into thermal groups is maintained naturally depends on the relative strength of the cohesive and disruptive forces. Those radicals such as OH and CN in which the bonds are very strong act as single thermal groups under all ordinary conditions. Those with somewhat weaker bonding— CO_3 , SO_4 , NO_3 , etc.,—also act as single units at the lower temperatures. Thus we find that at the initial points of both the first and second segments of the specific heat curves there are two groups in MnCO_3 , three in Na_2CO_3 , four in $\text{KAl}(\text{SO}_4)_2$, five in $\text{Ca}_3(\text{PO}_4)_2$, and so on. At higher temperatures, however, radicals of this class split up into two or more thermal groups. Still weaker radicals such as ClO_4 constitute two thermal groups even at low temperatures.

It was mentioned in Volume I that the boundary line between radicals and groups of independent atoms is rather indefinite. In general, the margin of bond strength required for a structural radical is relatively large, and we find many groups commonly recognized as radicals crystallizing in structures such as the CaTiO_3 cube in which the radical, as such, plays no part. The margin required in thermal motion is much smaller, particularly at the lower temperatures, and there are many atomic groups that act thermally in the same manner as the recognized radicals. In Li_2CO_3 , for example, the two lithium atoms act as a single thermal group, and the specific heat curve of this compound is similar to that of MgCO_3 rather than of Na_2CO_3 .

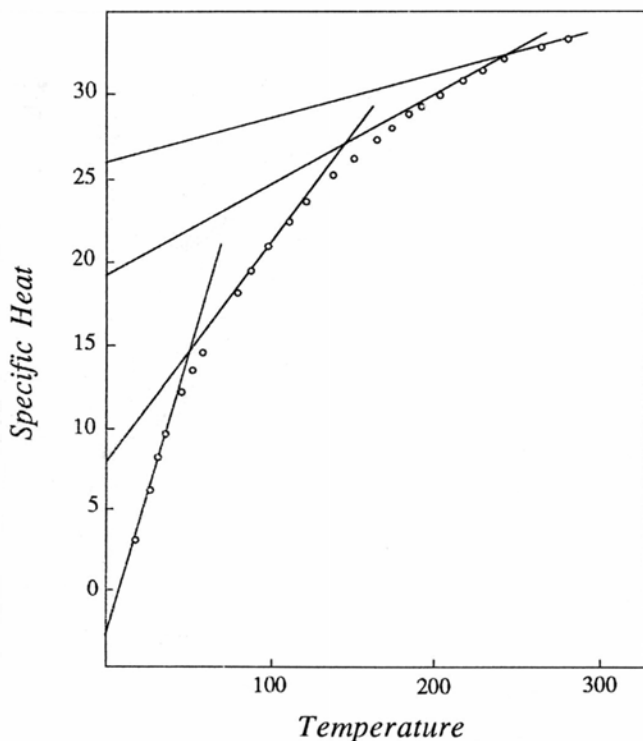
Extension of the thermal motion by breaking some of the stronger bonds at the higher temperatures gives rise to a variety of modifications of the specific heat curves. For example, MoS_2 has only two thermal groups in the lower range, but the S_2 combination breaks up as the temperature rises, and all atoms begin vibrating independently. VCl_2 similarly goes from one group to three. Splitting of the radical accounts for a change from two groups to three in SrCO_3 , from one to three in AgNO_3 , and from two to six in $(\text{NH}_4)_2\text{SO}_4$. All of these alterations take place at or prior to the first transitional point. Other compounds make the first transition on the initial basis and break up into more thermal groups later. In a common pattern, a radical that acts as one thermal group at the low temperatures splits into two groups in the temperature range of the second segment of the curve, just as the CO_3 radical in SrCO_3 , PbCO_3 and similar compounds does at a lower level. There are a number of structures such as KMnO_4 and KIO_3 , where this increases the number of groups in the molecule from two to three. $\text{Pb}_3(\text{PO}_4)_2$, in which there are two radicals, goes from five to seven, and so on.

The effect of water on crystallization is variable, depending on the strength of the cohesion. For example, $\text{BaCl}_2 \cdot 2\text{H}_2\text{O}$ acts as three thermal groups at the lower temperatures, the water molecules being firmly bound to the atoms of the compound. As the temperature increases these bonds give way, and the molecule begins vibrating on a five-group basis. In $\text{Al}_2(\text{SO}_4)_3 \cdot 6\text{H}_2\text{O}$ and in $\text{NH}_4\text{Al}(\text{SO}_4)_2 \cdot 12\text{H}_2\text{O}$ the bonds with the water molecules remain fixed through the entire experimental range, up to about 300°K , and the thermal groups in these hydrates are five and six respectively, just as in the corresponding anhydrous compounds.

An example of a drastic change in thermal behavior due to the disruption of inter-atomic bonds by thermal forces is shown in Fig. 12. The radical CrO_3 in the compound AgCrO_3 is a single thermal group at very low temperatures. There is a gradual separation into two groups in the temperature range up to the first transition point, and the change to the two-unit vibration is made on the basis of a two-group radical. At about 150K all four atoms in the radical begin vibrating independently, and the molecule undergoes a transition from the second segment of a three-group curve to the second segment of a

five-group curve. At about 250° K the compound makes the normal transition to three-unit vibration, continuing as five thermal groups.

FIGURE 12



The compounds used as examples in the foregoing discussion were selected mainly on the basis of the availability of experimental data within the significant temperature ranges. For an accurate definition of the slope of each of the straight line segments of any empirical curve it is necessary to have measurements in the temperature range where the deviations due to the proximity of a transition point are negligible. The examples have been taken from among those of the experimental results that satisfy this requirement.

In a theoretical treatment of specific heat such as that in this present work it is necessary to deal with this quantity on a per molecule basis. For practical application, however, it is more convenient to use the specific heat per unit of mass, and most of the collected data are expressed in this manner. It should be noted that the effect

of association into thermal groups is to reduce the specific heat per unit of mass. For this reason, the specific heat of most complex compounds is relatively low at low temperatures, and rises toward the values applicable to individual atoms as increasing temperature breaks up the original thermal groups.

The simplest organic compounds, those composed of only two or three structural units, generally divide into no more than two thermal groups. Many of the somewhat larger organic molecules, particularly among the ring structures and branched compounds, follow the same rule. The specific heat relations of these compounds are similar to those of the inorganic compounds, except that there are more organic compounds in which the thermal motion is restricted to one rotational unit. These substances, the hydrocarbons and some other compounds of the lower elements, undergo a transition to a positive initial level on reaching their first (and only) transition point. The resulting specific heat curve, the one illustrated in Fig. 3, is not much more than a straight line with a bend in it. A few compounds, including ethane and carbon monoxide, even omit the bend, and do not make the transition to the positive initial level.

Further addition of structural units, such as CH_2 groups, to the simple organic compounds results in the activation of *internal thermal groups*, units that vibrate thermally *within* the molecules. The general nature of the thermal motion of these internal groups is identical with that of the thermal motion of the molecule as a whole. But the internal motion is independent of the molecular thermal motion, and its scalar direction (inward or outward) is independent of the scalar direction of the molecular motion. Outward internal motion is thus coincident with outward molecular motion during only one quarter of the vibrational cycle. Since the effective magnitude of the thermal motion, which determines the specific heat, is the scalar sum of the internal and molecular components, each unit of internal motion adds one-half unit of specific heat during half of the molecular cycle. It has no thermal effect during the other half of the cycle when the molecule as a whole is moving inward.

Because of the great diversity of the organic compounds the specific heat patterns occur in a variety that is correspondingly large. The effect of internal motion in those of the organic compounds in which

it is present is well illustrated, however, by the specific heats of the normal paraffins. The values of the initial levels and the specific heat at T_1 for the compounds of this series in the range from C_3 (propane) to C_{16} (hexadecane) are listed in Table 21, together with the number of internal thermal units in the molecule of each compound.

TABLE 21: *SPECIFIC HEATS-PARAFFIN HYDROCARBONS*

	Internal Thermal Units	Initial Levels		Specific Heat at T_1
		1 st	2 nd	
Propane	0	-2.64	2.64	9.27
Butane	0	-2.64	2.64	9.27
Pentane	2	-2.64	6.62	13.90
Hexane	3	-2.64	7.95	16.22
Heptane	4	-3.96	9.27	18.54
Octane	5	-3.96	10.59	20.86
Nonane	6	-5.30	11.92	23.18
Decane	7	-5.30	13.24	25.49
Hendecane	8	-5.30	14.57	27.81
Dodecane	9	-6.62	15.89	30.12
Tridecane	10	-6.62	17.95	32.44
Tetradecane	11	-6.62	18.54	34.76
Pentadecane	12	-6.62	19.86	37.08
Hexadecane	13	-7.95	21.19	39.39

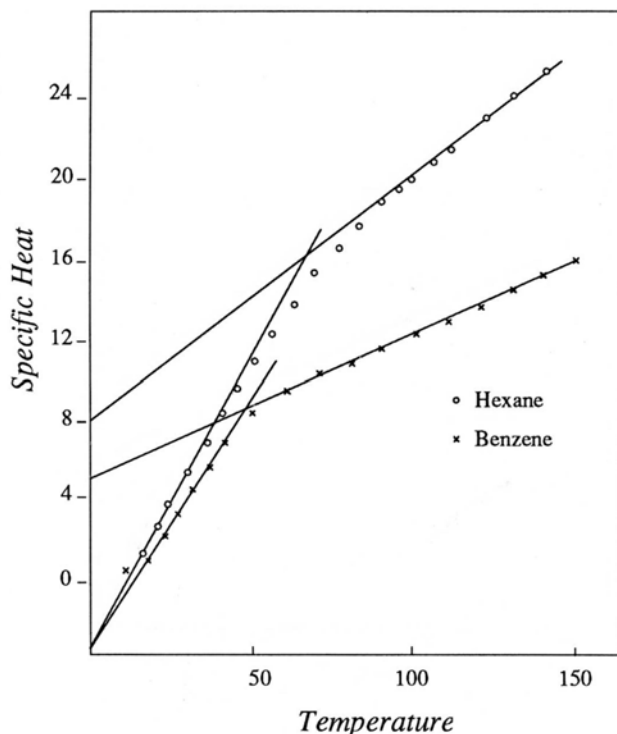
Propane and butane have only the two molecular thermal groups corresponding to the positive and negative ends of the molecules, and their specific heat at T_1 is the normal two-group value: 9.27. Beginning with two internal groups in pentane, each added CH_2 structural group becomes an internal thermal unit, and adds 2.317 to the total specific heat of the molecule at the transition point. The initial level of the first segment of the specific heat curve is -2.64 (the two-group value) in the lower compounds, and changes slowly, adding units of -1.32, as the length of the chain increases. The initial level of the second segment is 2.64 in butane and propane. In the higher compounds, each of which consists of n structural groups (CH_2 and CH_3), this second initial level is $1.324 n$.

The values thus derived theoretically are all consistent with the experimental curves. In a few cases the intersection of the two curve

segments may not coincide with the calculated specific heat of the transition point, but these deviations, if they are real, are small enough to be explainable on the basis of changes in the temperature factors, the nature of which will be one of the subjects of discussion in Chapter 7.

Branching of a hydrocarbon chain tightens the structure and tends to reduce the number of internal thermal units. For example, octane has five internal thermal units, and a specific heat of 20.86 at the transition point. But 2,2,4-trimethyl pentane, a branched compound with the same composition, has no internal motion at all, and the T_1 specific heat of this compound is 9.27, identical with that of the C_3 paraffin, propane. Ring formation has a similar effect. Ethyl-benzene and the xylenes, which are also C_8 compounds, have some internal motion, but their T_1 specific heats are 11.59 (one internal unit) and 13.90 (two internal units) respectively, well below the octane level. In Fig. 13 the specific heat curves of hexane (straight chain) and benzene (ring), both C_6 hydrocarbons, are contrasted.

FIGURE 13: *SPECIFIC HEAT*



The subject matter of this and the preceding five chapters consists of various aspects of the volumetric and thermal relations of material substances. The study of these relations was the principal avenue of approach to the clarification of basic physical processes that ultimately led to the identification of the physical universe as a universe of motion, and the determination of the nature of the fundamental features of that universe. These relations were examined in great detail over a period of many years, during which thousands of experimental results were analyzed and studied. Incorporation of the accumulated mass of information into the theoretical structure was the first task undertaken after the formulation of the postulates of the Reciprocal System of theory, and it has therefore been possible to present a reasonably complete description of each of the phenomena thus far discussed, including what we may call the small-scale effects.

Beginning with the next chapter, we will be dealing with subjects not covered in the inductive phase of the theoretical development. In this second phase, the deductive development, we are extending the application of the theory to all of the other major fields of physical science, in order to demonstrate that it is, in fact, a *general* physical theory. Obviously, where the area to be covered is so large, no individual investigator can expect to carry the development into great detail. Consequently, some of the conclusions expressed in the subsequent pages with respect to the small-scale features of the areas covered are subject to a degree of uncertainty. In other cases it will be necessary to leave the entire small-scale pattern for some future investigation.

CHAPTER 7

Temperature Relations

As explained in introducing the comparisons of the theoretical specific heats with experimental results, the curves in Fig. 5 to 13 verify only the specific heat pattern, the temperature scale of each curve being adjusted to the empirical results. In order to complete the definition of the curves we will now turn our attention to the temperature relations.

All of the distinctive properties of the different kinds of matter are determined by the rotational displacements of the atoms of which these substances are composed, and by the way in which the displacements enter into the various physical phenomena. As stated in Volume I,

The behavior characteristics, or *properties*, of the elements are functions of their respective displacements. Some are related to the total net effective displacement... some are related to the electric displacement, others to the magnetic displacement, while still others follow a more complex pattern. For instance, valence, or chemical combining power, is determined by either the electric displacement or one of the magnetic displacements, while the inter-atomic distance is affected by *both* the electric and magnetic displacement, but in different ways.

The great variety of physical phenomena, and the many different ways in which different substances participate in these phenomena result from the extension of this “more complex pattern” of behavior to a still greater degree of complexity. One of these more complex patterns was examined in Chapter 4, where we found that the response of the solid structure to compression is related to the cross-section against which the pressure is exerted. The numerical magnitude involved in this relation is determined by the product of the effective cross-sectional factors, together with the number of rotational units that participate in the action, a magnitude that determines the force

per unit of the cross-section. Inasmuch as one of the dimensions of the cross-section may take either the effective magnetic displacement, represented by the symbol b in the earlier discussion, or the electric displacement, represented by the symbol c , two new symbols were introduced for purposes of the compressibility chapter: the symbol z to represent the second displacement entering into the cross-section (either b or c), and the symbol y to represent the number of effective rotational units (related to the third of the displacements). The a - b - c factors were thus represented in the form a - z - y .

The values of these factors relative to the positions of the elements in the periodic table follow the same general pattern in application to specific heat as in compressibility, and most of the individual values are either close to those applying to compressibility or systematically related to those values. We will therefore retain the a - z - y symbols as a means of emphasizing the similarity. But the nature of the thermal relations is quite different from that of the relations that apply to compressibility. The temperature is not related to a cross-section; it is determined by the total effective rotation. Consequently, instead of the *product*, azy , of the effective rotational factors, the numerical magnitude defining the temperature scale of the thermal relations is the *scalar sum*, $a+z+y$, of these rotational values.

This kind of a quantity is quite foreign to conventional physics. The scalar aspect of vectorial motion is recognized; that is, speed is distinguished from velocity. But orthodox physical thought does not recognize the existence of motion that is *inherently* scalar. In the universe of motion defined by the postulates of the Reciprocal System of theory, on the other hand, all of the basic motions are inherently scalar. Vectorial motions can exist only as additions to certain kinds of combinations of the basic scalar motions.

Scalar motion in one dimension, when seen in the context of a stationary spatial reference system, has many properties in common with vectorial motion. This no doubt accounts for the failure of previous investigators to recognize its existence. But when motion extends into more than one dimension there are major differences in the way these two types of motion present themselves (or do not present themselves) to observation. Any number of separate vectorial motions of a point can be combined into a single resultant, and the

position of the point at any specified time can be represented in a spatial system of reference. This is a necessary consequence of the fact that vectorial motion is motion *relative to that system of reference*. But scalar motions cannot be combined vectorially. The resultant of scalar motion in more than one dimension is a scalar sum, and it cannot be identified with any *one* point in spatial coordinates. Such motion is therefore incapable of representation in a spatial reference system of the conventional type. It does not follow, however, that inability to represent this motion within the context of the severely limited kind of reference system that we are accustomed to use means that such motion is non-existent. To be sure, our direct perception of physical events is limited to those that can be represented in this type of a reference system, but Nature is not under any obligation to stay within the perceptive capabilities of the human race.

As pointed out in Chapter 3, Volume I, where the subject of reference systems was discussed at length, there are many aspects of physical existence (that is, many motions, combinations of motions, or relations between motions) that cannot be represented in *any* single reference system. This is not, in itself, a new, or unorthodox conclusion. Most modern physicists, including all of the leading theorists, have realized that they cannot accommodate all of present-day physical knowledge within the limitations of fixed spatial reference systems. But their response has been the drastic step of cutting loose from physical reality, and building their fundamental theories in a shadow realm where they are free from the constraints of the real world. Heisenberg states their position explicitly. "The idea of an objective real world whose smallest parts exists objectively in the same sense as stones and trees exist, independently of whether or not we observe them... is impossible,"⁸ he says. In the strange half-world of modern physical theory the only realities are mathematical symbols. Even the atom itself is "in a way only a symbol,"⁹ Heisenberg tells us. Nor is it required that symbols be logically related or understandable. Nature, these front rank theorists contend, is inherently ambiguous and subject to uncertainties of a fundamental and inescapable nature. "The world is not intrinsically reasonable or understandable," Bridgman explains, "It acquires these properties in ever-increasing degree as we ascend from the realm of the very little to the realm of everyday things."¹⁰

What the Reciprocal System of theory has done in this area is to show that once the true status of the physical universe as a universe of motion is recognized, and the properties of space and time are defined accordingly, there is no need for the retreat from reality, or for the attempt to blame Nature for the prevailing inability to understand the basic relations. The existence of phenomena not capable of representation in a spatial reference system is a fact that we must come to terms with, but the contribution of the Reciprocal System has been to show that the phenomena outside the scope of the conventional spatial reference systems can be described and evaluated in terms of the *same real entities* that exist within the reference system. The scalar sum of the magnitudes of motions in different dimensions, the quantity that we will now use in analyzing the temperature relations, is an item of this nature. It is just as real as any other physical quantity, and its components, the motions in the individual dimensions, are motions of the same nature as those one-dimensional scalar motions that *are* capable of representation in the spatial reference systems, even though the scalar sum cannot be so represented in any manner accessible to our direct perception.

In the theoretical minimum situation, where the effective thermal factors are 1-0-0, and the scalar sum of these factors is one unit, the temperature of the initial negative level is one unit out of the total of 128 that corresponds to the full 510.7 degrees temperature unit of the condensed states. But since the thermal motion is effective in only one direction, the ratio becomes $\frac{1}{256}$, and the zero point temperature, T_0 , the temperature at which the thermal motion counterbalances the negative initial level of vibration, is 1.995 K. For a substance with thermal factors a, z, and y, and the normal $\frac{2}{9}$ initial specific heat level, we then have

$$T_0 = 1.995 (a+z+y) \text{ degrees K} \quad (7-1)$$

This value completes the definition of the specific heat curves by defining the temperature scales. It will be more convenient, however, to work with another of the fixed points on the curves, the first transition point, T_1 . As this is the unit specific heat level on the initial linear section of the curve, while T_0 is $\frac{2}{9}$ unit above the initial point, the temperature of the first transition point is

$$T_1 = 8.98 (a+z+y) \text{ degrees K} \quad (7-2)$$

Thermal factors of the elements for which reliable specific heat patterns are available, and the corresponding theoretical first transition temperatures (T_1) are listed in Table 22, together with the T_1 values derived from curves of the type illustrated in Figs. 5 to 13, in which the temperature scale is empirical. In effect, this is a comparison between theoretical and experimental values of the temperature scales of the specific heat curves. The experimental values are subject to some uncertainty, as they have been obtained by inspection from graphs in which the linear portions of the curves were also drawn from visual inspection. Greater accuracy could be attained by using more sophisticated techniques, but the time and effort required for this refinement did not appear to be justified for the purposes of this initial investigation of the subject.

The compressibility factors derived in Chapter 4, with a few values restated in different, but equivalent, terms, are shown in the table for comparison with the corresponding thermal factors. The principal determinants of the compressibility values, aside from the effect of the pressure level itself (including the internal pressure), were found to be the magnitude and sign (positive or negative) of the displacement in the electric dimension. The rotational group to which the element belongs (determined by the magnetic displacements) is much less significant. In the thermal situation the rotational group becomes the dominant influence. The elements of Group 3B (magnetic displacements 3-3), midway in the group order, generally have thermal factors close to the compression values. In half of the 3B elements included in the table the deviation is no more than one unit. But in each direction from this central group there is a systematic deviation from the compressibility values, upward in the lower groups and downward in the higher groups. Every element *above* number 42, molybdenum, that is included in the table, with one exception, has thermal factors either equal to or less than the corresponding compressibility factors. Every element *below* molybdenum, with three exceptions (two of which are alkali metals), has thermal factors that are either equal to or greater than the corresponding compressibility factors.

TABLE 22: *EFFECTIVE ROTATIONAL FACTORS*

	Factors			T ₁				Factors			T ₁		
	Comp	Therm.	n	Tot.	Calc.	Obs.		Comp.	Therm.	Tot.	Calc.	Obs.	
Li	4-1-1	4-2-1	2	14	126	131	Y	4-2-1	4-3-1	8	72	71	
Be		4-1-1	2	12	108	110	Zr	4-8-1	4-4-1	9	81	84	
	4-4-1	4-2-1	2	14	314	323	Mo	4-8-2	4-8-2	14	126	129	
			8	56	314	323			4-6-2	12	108	107	
		4-1-1	2	12	269	267	Ru	4-8-2	4-8-2	14	126	128	
			8	48	269	267			4-6-2	12	108	107	
B		4-1-1	8	48	431	420	Rh	4-8-2	4-8-1	13	117	117	
C-d	4-6-1	4-4-1	8	72	647	635			4-6-1	11	99	95	
C-g	4-2-1	4-3-1	8	64	575	578	Pd	4-6-2	4-4-2	10	90	91	
Na	4-1-1	4-1-1		6	54	52			4-4-1	9	81	78	
Mg	4-4-1	4-1-1	2	12	108	109	Ag	4-4-2	4-3-1	8	72	72	
Al		3-1-1	2	10	90	91	Cd	4-4-1	2-2-1	5	45	46	
	4-5-1	4-2-1	2	14	126	131	In	4-4-1	4-6-2	12	108	105	
Si		4-1-1	2	12	108	112	Sn	4-4-1	4-2-1	7	63	66	
	4-4-1	4-6-2	2	24	216	220			4-1-1	6	54	57	
P-r		4-6-2	2	24	216	207	Sb	4-4-1	4-3-1	8	72	68	
P-w		4-2-1		7	63	66	Te	4-3-1	4-2-1	7	63	61	
S	4-1-1	4-4-1		9	81	84	I		2-2-1	5	45	44	
Cl		4-2-1		7	63	62	Xe		1-1-0	2	18	19	
Ar		1-1-1		3	27	28	Cs	4-1-1	1-1-0	2	18	17	
K	4-1-1	2-1-1		4	36	32	Ba	4-2-1	2-1-1	4	36	34	
Ca	4-3-1	4-3-1		8	72	76	La	4-4-1	2-2-1	5	45	42	
Sc		4-6-1		11	99	103	Pr	4-4-1	1-1-1	3	27	27	
Ti		4-5-1		10	90	88	Nd	4-4-1	1-1-1	3	27	31	
	4-8-1	4-8-2		14	126	124	Sm	4-4-1	2-1-1	4	36	36	
V	4-8-1	4-8-3		15	135	133	Eu	4-4-1	2-1-1	4	36	33	
Cr		4-6-2		12	108	107	Gd	4-4-1	2-2-1	5	45	48	
	4-8-1					162	Tb	4-4-1	2-2-1	5	45	44	
Mn		4-8-2		14	126	128	Dy	4-4-1	2-2-1	5	45	41	
	4-8-1	4-8-1		13	117	115	Ho	4-4-1	2-1-1	4	36	33	
Fe		4-5-1		10	90	92	Er	4-4-1	1-1-1	3	27	28	
	4-8-1	4-8-4		16	144	142	Tm	4-4-1	1-1-1	3	27	29	
Co		4-6-2		12	108	108	Yb	4-2-1	2-1-1	4	36	37	
	4-8-1	4-8-2		14	126	126	Hf		4-3-1	8	72	71	
Ni		4-6-1		11	99	100	Ta	4-8-2	4-3-1	8	72	74	
	4-8-1	4-8-2		14	126	131	W	4-8-3	4-6-2	12	108	108	
Cu		4-6-1		11	99	97	Re		4-4-2	10	90	93	
	4-6-1	4-6-2		12	108	108			4-4-1	9	81	78	
Zn	4-4-1	4-3-1		8	72	73	Ir	4-8-3	4-6-1	11	99	98	
Ga		2-1-1		4	36	36			4-5-1	10	90	88	
Ge	4-4-1	4-8-1		13	117	119	Pt	4-8-2	4-3-1	8	72	76	

	Factors			T ₁				Factors			T ₁		
	Comp	Therm.	n	Tot.	Calc.	Obs.		Comp.	Therm.	Tot.	Calc.	Obs.	
As	4-4-1	4-6-2		12	108	106	Au	4-6-2	4-1-1	6	54	57	
Se	4-1-1	4-3-1		8	72	75	Hg		2-1-1	4	36	32	
Br		4-2-1		6	56	54	Tl	4-4-1	2-1-1	4	36	34	
Kr		1-1-0		2	18	20	Pb	4-4-1	2-1-1	4	36	33	
Rb	4-1-1	1-1-0		2	18	20	Bi	4-3-1	2-2-1	5	45	44	

It was noted in Chapter 4 that in compression the lowest electropositive elements do not take the minimum 1-1-1 factors of their electronegative counterparts, but have a=4 in all of the elements of this class investigated by Bridgman. The reason for this difference in behavior is not yet known (although it is no doubt connected with the all-positive nature of the rotational displacement of these elements), but it is even more pronounced in the thermal factors. Except for the alkali metals above sodium, which, as noted above, have thermal factors even lower than the compressibility values, the lower electropositive elements not only maintain the 6-unit minimum (4-1-1 or equivalent) but raise the effective magnitudes of their thermal factors still farther by omitting the n=1 section of the specific heat curve based on equation 5-6, and going immediately to n=2, which increases the temperature scale by a factor of 8. This pattern is followed by boron and carbon, and in part, by beryllium. The corresponding members of the next higher group, magnesium, aluminum, and silicon, also have n=2 from the start of the thermal motion, but here the second unit is one-dimensional rather than three-dimensional. Beryllium combines the two patterns. It has the same thermal factors as lithium, but a dimensional multiplier halfway between those of lithium and boron, the two adjoining elements.

The option of one dimension or three dimensions is open whenever motion advances from one unit to two units, but not under any other conditions. Three-dimensional motion of one displacement unit is meaningless, as $1^3 = 1$. After two units there is no option, as there cannot be more than two units in linear succession, for reasons that were discussed in Volume I. But two-unit motion can be either one-dimensional or three-dimensional. At the point where the advance from one to two units takes place, the motion is therefore able to take the dimensions that are best suited to the existing situation.

A one-dimensional increase in the value of n results in increasing the temperature scale by a factor of 2 rather than 8. The alkali metals, which diverge from the normal electropositive behavior in a number of respects because of their low electric displacement, follow the same pattern as the elements listed in the preceding paragraph, but one step lower, as indicated in the following comparison:

Group	Alkalis	Other Positive
1B	$n=2$	$n=8$
2A	$4-x-x$	$n=2$
2B	$1-1-x$	$4-x-x$

As we found in the specific heat investigation, the electronegative elements below displacement 7 have a half-size initial negative specific heat level: $\frac{1}{9}$ unit instead of the normal $\frac{2}{9}$. It might be expected that this would result in a net effective specific heat of $\frac{8}{9}$ unit or $2\frac{2}{3} R$, at the transition point instead of the $\frac{7}{9}$ unit ($2\frac{1}{3}R$) that exists when the initial negative level is $\frac{2}{9}$ unit. But it is quite clear from the measured specific heat values that this is not true. The first transition point in the specific heat curves of the electronegative elements is $2\frac{1}{3} R$ just as it is in the curves with the $\frac{2}{9}$ unit ($\frac{2}{3}R$) negative initial level. Apparently the restriction that prevents the existence of the more negative initial level in the specific heat of these elements is gradually eliminated as the temperature rises, so that at the transition point the effective negative component of the specific heat is the normal $\frac{2}{9}$ unit.

The thermal factors of the higher inert gases, krypton and xenon, which have no rotation in the electric dimension, are 1-0-0 rather than 1-1-0, as in compressibility. This is a peculiarity of the mathematics, and has no physical significance. In both cases the meaning of the symbols is that the effective magnitude is determined entirely by the factors a and z . In multiplication this requires a unit value in the y position, whereas in addition a zero is required for the same purpose. But this equivalence of the 1-1-1 compressibility and 1-1-0 thermal factors does not mean that 1-1-1 *thermal* and 1-1-0 thermal are equivalent. The 1-1-1 thermal combination is the minimum for a substance with effective rotational displacement in all three

dimensions. Where the thermal factors drop to 1-1-0, as indicated for rubidium and cesium, there is no effective displacement in the electric dimension, and the thermal motion is following the inert gas pattern. Such behavior is uncommon, but it is not without precedent in other properties. We found in Chapter 1, for instance, that a number of elements, including the halogens, the elements corresponding to the alkalis on the opposite side of the inert gases, have inter-atomic distances in one or two dimensions that are similarly based on magnetic rotation only.

Since the empirical values listed in Table 22 are subject to a considerable degree of uncertainty, small differences between them and the calculated values have no significance. In some cases, however, the discrepancy is large enough to be real, and further study of the thermal relations of these elements will be required. Only one of the experimental values shown in the table, one of those applicable to chromium, is too far from any theoretical temperature to be capable of explanation on the basis of the theoretical information now available.

As brought out in the discussion of the general pattern of the specific heat curves in Chapter 5, in many substances there is a change in the temperature scale of the curve at the first transition point (T_1), as a result of which the first and second segments of the curve do not intersect at the $2\frac{1}{3}R$ end point of the lower segment of the curve in the normal manner. This change in scale is due to a transition to the second set of thermal factors given, for the elements in which it occurs, in Table 22. With the benefit of the information that we have developed regarding the factors that determine the temperature scale we can now examine the quantitative aspects of these changes.

As an example, let us look at the specific heat curve of molybdenum, Fig. 9, which, as previously noted, also applies to ruthenium. The thermal factors applicable to these elements at low temperatures are 4-8-2, identical with the compressibility factors. The first transition point, specific heat 4.63, is reached at 126K on the basis of these factors. The corresponding empirical temperatures, determined by inspection of the trend of the experimental values of the specific heats, are 129 for molybdenum and 128 for ruthenium, well within the range of uncertainty of the techniques employed in estimating the

empirical values. If the thermal factors remained constant, as they do in the “regular” pattern followed by such elements as silver, Fig. 5, there should be a transition to $n=2$ at this 126 K temperature, and the specific heat above this point would follow the extension of a line from the initial level of 3.89 to 4.63 at 126 K. But instead of continuing on the 4-8-2 basis, the thermal factors decrease to 4-6-2 at the transition point. These factors correspond to a transition temperature of 108 K. The specific heat of the molecule therefore undergoes an isothermal increase at 126 K to the extension of a line from the initial level of 3.89 to 4.63 at 108 K, and follows this line at higher temperatures. The effect of the isothermal increase in the specific heat of the individual molecules is, of course, spread out over a substantial temperature range in application to a solid aggregate by the distribution of molecular velocities.

The temperature of the subsequent transition points and the end points of the various segments of the specific heat curves can be calculated from the temperatures of the first transition points by applying the relative values listed in Chapter 5 to the appropriate values of T_1 . An approximate agreement between the empirical data and the higher transition points thus calculated is indicated, but the angles at which the upper segments of the curves intersect are too small to permit any close empirical definition of the temperature of intersection. The only one of the end points that has any real significance is the end point of the last segment of the curve applicable to the substance under consideration. This is the temperature limit of the solid. Any further addition of heat initiates the transition to the liquid state.

Inasmuch as it is the individual molecule that reaches its thermal limit at the solid end point, it is the individual molecule that makes the transition to the liquid state. Physical state is thus basically a property of the individual molecule rather than a property of the aggregate, as seen in conventional physical theory. The state of the aggregate is merely a reflection of the state of the majority of its constituents. Recognition of this fact some forty years ago, in the early stages of the investigation that led to the results now being reported, was a major step in the clarification of physical fundamentals that ultimately opened the door to the formulation of a general physical theory.

The liquid state has long been an enigma to conventional physics. As expressed by V. F. Weisskopf, "A liquid is a highly complex phenomenon in which the molecules stay together yet move along each other; it is by no means obvious why such a strange substance should exist."¹¹ Weisskopf goes on to speculate as to what the outcome would be if physicists knew the fundamental principles on which atomic structure is based, as present-day theory sees them, but "never had any occasion to see structures in nature." He doubts if these theorists would ever be able to predict the existence of liquids.

In the Reciprocal System of theory, on the other hand, the liquid state is a necessity, an intermediate condition that must necessarily exist between the solid and gaseous states. When the thermal motion of a molecule reaches equality with the inward progression of the natural reference system in one dimension of the region outside unit distance, the cohesive force in that dimension is eliminated. The molecule is then free to move in that dimension, while it is held in a fixed position, or a fixed average position, in the other dimensions by the cohesive forces that are still operative. The temperature at which the freedom in one dimension is reached is the melting point of the aggregate, because any additional thermal energy supplied to the aggregate is absorbed in changing the state of additional molecules until the remaining content of solid molecules reaches the percentage that can be accommodated within the liquid aggregate.

These remaining solid molecules are gradually converted to the liquid state in a temperature range above the melting point. Thus the liquid aggregate in this range contains a percentage of solid molecules, while the solid aggregate in a similar temperature range below the melting point contains a percentage of liquid molecules. The presence of these "foreign" molecules has a significant effect on the physical properties of matter in both of these temperature ranges, an effect which, as we will see in the subsequent discussion of the liquid state, can be evaluated accurately by using probability relations to determine the exact proportions in which molecules of the two states exist at each temperature.

While the end point of the solid state is the temperature at which the inter-molecular forces reach an equilibrium at the unit level, arrival at this end point does not mean automatic entry into the liquid state.

It merely means that the cohesive forces of the solid are no longer operative in all three dimensions, and therefore do not *prevent* the free movement in one dimension of space that is the distinguishing characteristic of the liquid state. In order to actually become a liquid the molecule must also conform to all of the essential requirements of the liquid state. The significant point here is that a liquid *molecule* is limited to certain specific temperatures. A liquid *aggregate* can take any temperature within the liquid range, but only because the aggregate temperature is an average of a large number of the restricted individual values.

This same restriction to one of a limited set of values also applies to the temperature of the solid molecule, but in the vicinity of the melting point the solid is at a high time region temperature level, where the proportionate change from one possible value, n units, to the next, $n + 1$ units, is small. The motion of the liquid state, on the other hand, is in the region outside unit space, and is equivalent to gas motion in the one dimension in which the thermal energy exceeds the solid state limit. As we saw in Chapter 5, temperatures in the vicinity of the melting point are very low on the scale applicable to this outside region, and the proportionate change from n to $n + 1$ is large. The intervals between the possible temperatures of liquid molecules are therefore large enough to be significant.

Because of the limitation of the liquid temperatures to specific values, the temperature at which a molecule qualifies as a liquid is not the end point temperature of the solid state, but a higher value that includes the increment necessary to bring the end point temperature up to the next available liquid level. This makes it impossible to calculate melting points from solid state theory alone. Such calculations will have to wait until the relevant liquid theory is developed in a subsequent volume in this series, or elsewhere. But the temperature increment beyond the solid end point is small compared to the end point temperature itself, and the end point is not much below the melting point. A few comparisons of end point and melting point temperatures will therefore serve to confirm, in a general way, the theoretical deductions as to the relation between these two magnitudes.

There is a considerable degree of uncertainty in the experimental results at the high temperatures reached by the melting points of

many of the elements, and there are also some theoretical aspects of the thermal situation in the vicinity of the melting point that have not yet been fully explored. The examples for discussion in this initial approach to the subject have been selected from among those in which these uncertain elements are at a minimum. First, let us look at element number 19, potassium. This element has a specific heat curve of the type identified by the notation $n=3$ in Fig. 4. Its thermal factors are 2-1-1, and it maintains the same factors throughout the entire solid range. As indicated in Chapter 5, the end point temperature of this type of curve is 9.32 times the temperature of the first transition point. This leads to an end point temperature of 336 K. The measured melting point is 337 K. In this case, then, the solid end point and the melting point happen to coincide within the limits of accuracy of the investigation.

Chlorine, an element only two steps lower in the atomic series than potassium, but a member of the next lower group, has the lower type of specific heat curve, with $n=2$. The end point temperature of this curve is 4.56 on the relative scale where the first transition point is unity. The thermal factors that determine the transition point, and are applicable to the first segment of the curve, are 4-2-1, but if these factors are applied to the end point they lead to an impossibly high temperature. It is thus apparent that the factors applicable to the second segment of the curve are lower than those applicable to the first segment, in line with the previously noted tendency toward a decrease in the thermal factors with increasing temperature. The indicated factors applicable to the end point in this case are the same 2-1-1 combination that we found in potassium. They correspond to an end point temperature of 164 K, just below the melting point at 170 K, as the theory requires.

Next we look at two curves of the $n=4$ type, the end point of which is at a relative temperature of 17.87. On the basis of thermal factors 4-6-1, the absolute temperature of the end point is 1765 K, which is consistent with the melting points of both cobalt (1768) and iron (1808). Here, too, the indicated factors at the end point are lower than those applicable to the first segment of the specific heat curve, but in this case there is independent evidence of the decrease. Cobalt, which has the factors 4-8-2 in the first segment is already down to

4-6-1 at the second transition point, while iron, the initial factors of which are also 4-8-2, has reached 4-6-2 at this point, with two more segments of the curve in which to make the additional reduction.

Compounds of elements above group 1B, or having a significant content of such elements, follow one or another of the Type 1 patterns that have been illustrated by examples from the elements. The hydrocarbons and other compounds of the lower group elements have specific heat curves of type 2 (Fig. 3) in which the end point is at a relative temperature of 1.80. As an example of this class we can take ethylene. The thermal factors of these lower group compounds are limited to 1-1-1, 2-1-1, and the combination value $1\frac{1}{2}$ -1-1. As we found in Volume I, however, the two groups of atoms of which ethylene and similar compounds are composed are inside one time region unit of distance. They therefore act jointly in thermal interchange rather than acting independently in the manner of two inorganic radicals, such as those in NH_4NO_3 . Each group contributes to the thermal factors of the molecule, and the value applicable to the molecule as a whole is the sum of the two components. Ethylene uses the 1-1-1 and $1\frac{1}{2}$ -1-1 combinations. A difference of this kind between the two halves of an organic molecule is quite common, and no doubt reflects the lack of symmetry between the positive and negative components that was the subject of comment in the discussion of organic structure. The combined factors amount to a total of $6\frac{1}{2}$ units. This corresponds to a transition point at 58 K, which agrees with the empirical curve, and an end point at 104 K, coincident with the observed melting point.

The joint action of the two ends of an organic molecule that combines their thermal factors in the temperature determination is maintained when additional structural units are introduced between the end groups. As brought out in Chapter 6, such an extension of the organic type of structure into chains or rings also results in the activation of additional thermal motions of an independent nature within the molecules. The general nature of this internal motion was explained in the previous discussion. The same considerations apply to the transition point temperature, except that the internal motion is independent of the molecular motion in vectorial direction as well as in scalar direction. It is therefore distributed three-dimensionally, and the fraction in the direction of the molecular motion is $\frac{1}{3}$ rather

than $\frac{1}{2}$. Each unit of internal motion thus adds $\frac{1}{8}$ of 8.98 degrees, or 1.12 degrees K to the transition point temperature. With the benefit of this information we are now able to compute the temperatures corresponding to the specific heats of the paraffin hydrocarbons of Table 21. These values are shown in Table 23.

TABLE 23:
TEMPERATURES OF CRITICAL POINTS—PARAFFIN HYDROCARBONS

<i>Thermal Factors</i>						
	Transition Point	Total		End Point	Total	
Propane	1-1-1	1-1-1	6	1-1-1	1-1-0	5
Butane	1-1-1	1-1- $\frac{1}{2}$	5 $\frac{1}{2}$	2-1-1	1 $\frac{1}{2}$ -1-1	7 $\frac{1}{2}$
Pentane	1 $\frac{1}{2}$ -1-1	1 $\frac{1}{2}$ -1-1	7	2-1-1	2-1-1	8
Hexane & above	2-1-1	1 $\frac{1}{2}$ -1-1	7 $\frac{1}{2}$	2-1-1	2-1-1	8
<i>Temperatures</i>						
	Internal Units	T ₁	End Point Internal	Factors Total	End Point	Melting Point
Propane	0	54	0	5	81	85
Butane	0	50	1	8 $\frac{1}{2}$	137	138
Pentane	2	65	1	9	145	143
Hexane	3	71	3	11	178	179
Heptane	4	72	3	11	178	182
Octane	5	73	5	13	210	216
Nonane	6	74	5	13	210	220
Decane	7	75	7	15	242	243
Hendecane	8	76	7	15	242	247
Dodecane	9	77	8	16	259	263
Tridecane	10	79	8	16	259	268
Tetradecane	11	80	9	17	275	279
Pentadecane	12	81	9	17	275	283
Hexadecane	12	82	10	18	291	291

The first section of this table traces the gradual increase in the thermal factors as the molecule makes the transition from a simple combination of two structural groups, with properties that are similar to those of inorganic binary compounds, except for the joint thermal

action due to the short inter-group distance, to a long-chain organic structure. The increase in the factors follows a fairly regular course in this range except in the case of butane. If the experimental values of the specific heat of this compound are accurate, its transition point factors drop back from the total of 6 that applies to propane to $5\frac{1}{2}$, whereas they would be expected to advance to $6\frac{1}{2}$. The reason for this anomaly is unknown. At the C_6 compound, hexane, the transition to the long-chain status is complete, and the thermal factors of the higher compounds as far as hexadecane (C_{16}), the limit of the present study, are the same as those of hexane.

In the second section of the table the transition point temperatures are calculated on the basis of 8.98 degrees K per molecular thermal factor, as shown in the upper section of the table, plus 1.12 degrees per effective unit of internal motion. The number of internal motions shown in Column 1 for each compound is taken from Table 21.

Columns 3 and 4 are the values entering into the calculation of the solid end point, Column 5. As the table indicates, some of the internal motions that exist in the molecule at the transition temperature are inactive at the end point. However, the active internal motion components are thermally equivalent to the molecular motions at this point, rather than having only $\frac{1}{8}$ of the molecular magnitude as they do at T_1 . This is a result of the general principle that the state of least energy takes precedence (in a low energy environment) in cases where alternatives exist. Below the transition point the internal thermal motions are necessarily one-dimensional. Above T_1 they are free to take either the one-dimensional or three-dimensional status. The energy at any given temperature above T_1 is less on the three-dimensional basis. This transition therefore takes place as soon as it *can*, which is at T_1 . At the melting point the energy requirement is greater after the transition to the liquid state. Consequently, this transition does not take place until it *must* because there is no alternative. A return to one-dimensional internal thermal motion is an available alternative that will delay the transition. This motion therefore gradually reverts back to the one-dimensional status, reducing the energy requirement, and the solid end point is not reached until all effective thermal factors are at the 8.98 temperature level. The end point temperature of Column 5 is then $8.98 \times 1.80 = 16.164$ times the total number of thermal factors shown in Column 4.

The calculated transition points are all in agreement with the empirical curves within the margin of uncertainty in the location of these curves. As can be seen by comparing the calculated solid end points with the melting points listed in the last column, the end point values are also within the range of deviation that is theoretically explainable on the basis of discrete values of the liquid temperatures. It is quite possible that there is some “fine structure” involved in the thermal relations of solid matter that has not been covered in this first systematic theoretical treatment of the subject. Aside from this possibility, it should be clear from the contents of this and the two preceding chapters that the theory derived by development of the consequences of the postulates of the Reciprocal System is a correct representation of the general aspects of the thermal behavior of matter.

CHAPTER 8

Thermal Expansion

As indicated earlier, addition of thermal motion displaces the inter-atomic equilibrium in the outward direction. A direct effect of the motion is thus an expansion of the solid structure. This direct and positive result is particularly interesting in view of the fact that previous theories have always been rather vague as to why such an expansion occurs. These theories visualize the thermal motion of a solid as an oscillation around an equilibrium position, but they fail to shed much light on the question as to why that equilibrium position should be displaced as the temperature rises. A typical “explanation” taken from a physics text says, “Since the average amplitude of vibration of the molecules increases with temperature, it seems reasonable that the average distance between the atoms should increase with temperature.” But it is not at all obvious why this should be “reasonable.” As a general proposition, an increase in the amplitude of a vibration does not, in itself, change the position of equilibrium.

Many discussions of the subject purport to supply an explanation by stating that the thermal motion is an anharmonic vibration. But this is not an explanation; it is merely a restatement of the problem. What is needed is a reason *why* the addition of thermal energy produces such an unusual result. This is what the Reciprocal System of theory supplies. According to this theory, the thermal motion is not an oscillation around a fixed average position; it is a simple harmonic motion in which the inward component is coincident with the progression of the natural reference system, and therefore has no physical effect. The outward component is physically effective, and displaces the atomic equilibrium in the outward direction.

From the theoretical standpoint, thermal expansion is a relatively unexplored area of physical science. Measurement of the expansion

of different substances at various temperatures is being pursued vigorously, and the volume of empirical data in this field is increasing quite rapidly. However, the practical effect of the *change* in the coefficient of expansion due to temperature variation is of little consequence, and for most purposes it can be disregarded. As stated in the physics text from which the “explanation” of the expansion was taken, “Accurate measurements do show a slight variation of the coefficient of expansion with the temperature. We shall ignore such variations.” This lack of significant practical application has limited the amount of theoretical attention that the subject has heretofore received. But one of the principal objectives of this present work is to demonstrate that the Reciprocal System is a *general* physical theory. However limited the practical use of the thermal expansion information may be, we will want to show that this expansion can be explained on the same basis as the other properties of matter, using the same principles and relations that are applied to those other properties, with only such modifications as are required by considerations peculiar to the expansion.

In general, the volumetric behavior of a solid in response to the application of heat is analogous to that of a confined gas. The volume of an unconfined gaseous aggregate is not capable of definition, but when this aggregate is confined a volumetric equilibrium is reached on a basis that is not essentially different from that which exists in the solid state, aside from those items which depend on whether the point of equilibrium is inside or outside unit distance. At constant pressure, the general gas equation (5-3), which describes the relation between the principal properties of the ideal gas, reduces to

$$V = kT \quad (8-1)$$

This is *Charles’ Law*. It tells us that at constant pressure the volume of an ideal gas (one that is entirely free from time region forces) is directly proportional to the absolute temperature.

The relation $E = PV$ (equation 4-3) is merely a restatement of the definition of pressure, in a different form, and is therefore valid in the time region (inside unit distance) as well as in the ideal gas state. Since $E = kT^2$ (equation 5-5) in the time region, it follows that in this region

$$PV = kT^2 \quad (8-2)$$

At constant pressure this reduces to

$$V = kT^2 \quad (8-3)$$

For the same reasons that were discussed in connection with the heat equations, the relation between volume and temperature expressed in equation 8-3 is valid in this simple form only where the thermal motion is limited to one rotational displacement unit. The general expression applicable to motion on any rotational basis is

$$V = \frac{kT^2}{n^3} \quad (8-4)$$

In our consideration of volume changes in solid structures due to the addition of thermal energy we will usually be interested mainly in the *coefficient of thermal expansion*, or derivative of volume with respect to temperature. This is obtained by differentiating equation 8-4.

$$\frac{dv}{dT} = \frac{2kT}{n^3} \quad (8-5)$$

Aside from the numerical constant k , this equation is identical with the specific heat equation 5-7, where the value of n in that equation is unity. Thus there is a close association between thermal expansion and specific heat. The correlation between the two quantities, in application to the elements, is essentially complete up to the first transition temperature defined in Chapter 5. For all of the elements on which sufficient data are available to enable locating the transition point, this transition temperature is the same for thermal expansion as it is for specific heat. Each element has a negative initial level of the expansion coefficient, the magnitude of which has the same relation to the magnitude at the transition point as in specific heat; that is, $\frac{2}{9}$ in most cases, and $\frac{1}{9}$ in some of the electronegative elements. It follows that if the coefficient of expansion at the transition point is equated to 4.63 specific heat, the first segment of the expansion curve is identical with the first segment of the specific heat curve.

Beyond the transition point the thermal expansion curve follows a course quite different from that of the specific heat, because of the difference in the nature of the two phenomena. Since the term n^3 is absent from the thermal expansion equation, the modification

of the expansion curve that takes place where motion of single units is succeeded by multi-unit motion involves a change in the coefficient k . The expansion is related to the effective energy (that is, to the temperature), irrespective of the relation between total energy and effective energy that determines the specific heat above the first transition point. The magnitude of the constant k that determines the slope of the upper segment of the expansion curve is determined primarily by the temperature of the end point of the solid state.

For purposes of this present discussion, the solid end point will be regarded as coincident with the melting point. As brought out in Chapter 7, this is, in fact, only an approximate coincidence. But the present examination of thermal expansion is limited to its general features. Evaluation of the exact quantitative relations will not be feasible until a more intensive study of the situation is undertaken, and even then it will be difficult to verify the theoretical results by comparison with empirical data because of the large uncertainties in the experimental values. Even the most reliable measurements of thermal expansion are subject to an estimated uncertainty of ± 3 percent, and the best available values for some elements are said to be good only to within ± 20 percent. However, most of the measurements on the more common elements are sufficiently accurate for present purposes, as all that we are here undertaking to do is to show that the empirically determined expansions agree, in general, with the theoretical pattern.

The total expansion from zero temperature to the solid end point is a fixed quantity, the magnitude of which is determined by the limitation of the solid state thermal motion (vibration) to the region within unit distance. At zero temperature the gravitational motion (outward in the time region) is in equilibrium with the inward progression of the natural reference system. The resulting volume is s_0^3 , the initial molecular volume. At the solid end point the thermal motion is also in equilibrium with the inward motion of the reference system progression, as this is the point at which the thermal motion is able to cross the time region boundary without assistance from gravitation. The thermal motion up to the end point of the solid state thus adds a volume equal to the initial volume of the molecule. Because of the dimensional situation, however, only a fraction of the added volume is effective in the region in which it is measured; that is, outside unit space.

For an understanding of the dimensional relations that are involved it is necessary to realize that all of the phenomena of the solid state take place inside unit space (distance), in what we have called the time region. The properties of motion in this region were discussed in detail at appropriate points in Volume I. This discussion will not be repeated here, but a brief review of the general situation, with particular reference to the dimensions of motion may be helpful. According to the fundamental postulates of the Reciprocal System, space exists only in association with time as motion, and motion exists only in discrete units. From this it follows that space and time likewise exist only in discrete units. Consequently, any two atoms that are separated by one unit of space cannot move any closer together in space, as this would require the existence of fractional units. These atoms may, however, accomplish the *equivalent* of moving closer together in space by moving outward in time. All motion in the time region, the region inside unit space, is motion of this kind: motion in time (equivalent space) rather than motion in actual space.

The first unit of thermal motion is a one-dimensional motion in time. At the transition point, T_1 , this motion has reached the full one-unit level. As already explained, only half of this unit is physically effective. One fully effective unit is required for escape from the time region, and the motion therefore enters a second time region unit. In this second unit a three-dimensional distribution of the motion is possible. But the motion in time that takes place in the time region has only a scalar connection with motion in the region outside unit space, which is motion in space. This is equivalent to a one-dimensional contact. Thus only one dimension of the three-dimensional time region motion is effective beyond the regional boundary. The effective fraction of the motion is $\frac{1}{8}$ of one unit, or $\frac{1}{16}$ of the total two-unit time region motion. The expansion is proportional to the effective component of the motion, and this means that the volumetric expansion from zero temperature to the solid end point, as measured in the region outside unit space, is also $\frac{1}{16}$, or 0.0625 of the initial volume. On a one-dimensional (linear) basis, this is 0.0208.

This is the relative expansion that would take place providing that no change in the volumetric determinants of the substance occurs above the reference temperature (usually room temperature). But such

changes occur more often than not, and, as has been explained, the volume changes accompanying an increase in temperature are normally in the direction of increased volume. The total expansion is 0.0625 of the initial volume corresponding to the volume at the solid end point. Where this theoretical initial volume is greater than the reference volume projected to zero temperature, the expansion expressed relative to the smaller volume is correspondingly increased. It follows that in most cases the linear expansion, as measured, is somewhat above 0.0208, generally in the range from this value up to about 0.028.

The increase in volume at the higher temperature, where it occurs, is generally due to structural rearrangements. The changes take place either in the inter-atomic distance, by reason of transitions from one of the types of orientation discussed in Chapter 1 to another, or in the crystal structure, or both. The expansion is related to the inter-atomic distance (s_0) rather than to the geometrical volume, and it is independent of the geometrical arrangement, but, as indicated in the preceding paragraph, a modification of the geometry does affect the relation of the volume at the solid end point to the reference volume at zero temperature.

In the NaCl type of structure the edge of the unit cube is equal to the inter-atomic distance. This cube contains one atom, and the ratio of the measured volume to what we may call the three-dimensional space, the cube of the inter-atomic distance, is therefore unity. In the body-centered cube the edge is $2/\sqrt{3}$ times the inter-atomic distance. Since the unit cube of this type contains two atoms, the ratio of volume to three-dimensional space is 0.770. The one-dimensional space, the edge of a hypothetical cube containing one atom, is then 0.9165 for the body-centered cube and 1.00 for the NaCl type structure. Transitions from one type of structure to the other modify the spatial relations accordingly. The values applicable to all five of the principal isometric crystal structures of the elements are listed in the following tabulation.

Face-centered cube	0.8909
Close-packed hexagonal	0.8909
Body-centered cube	0.9165
Simple (NaCl) cube	1.0000
Diamond (ZnS) cube	1.1547

The second segment of the thermal expansion curve has no negative initial level, because there is a positive expansion (that of the first segment) into which the initial level can extend. Like the transition from the liquid to the solid state, the transition from single units of motion to multi-unit motion involves a change in the zero datum applicable to temperature. The temperature T_0 , corresponding to the initial negative level, is eliminated, and the temperature of the end point, T_1 , of the first segment of the curve, which is $\frac{1}{2}T_0$ on this segment, is reduced to $\frac{1}{2}T_0$ on the second segment.

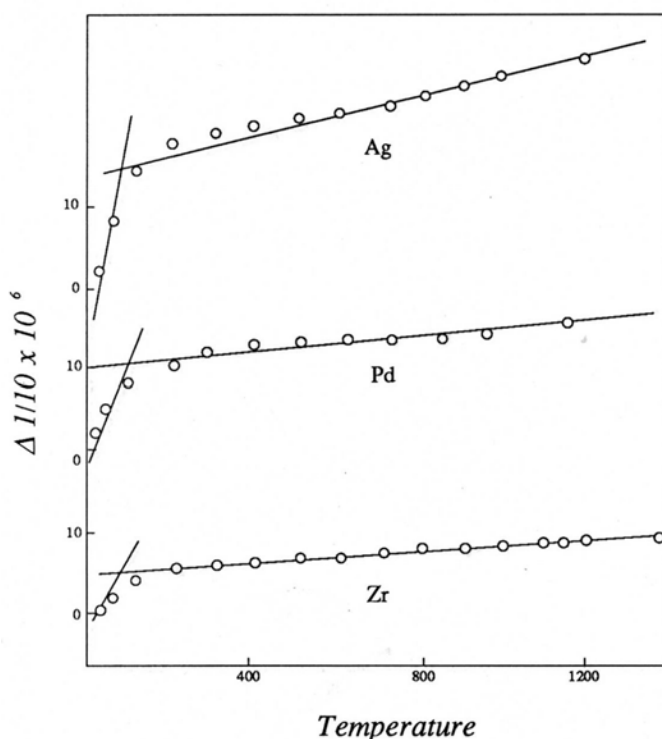
As brought out in Chapter 7, the minimum of the zero point temperature, T_0 , is equivalent to one of the 128 dimensional units that correspond to one full temperature unit, 510.7 degrees K. As the temperature rises, additional units of motion are activated, and the corresponding value when all 128 units are fully effective is thus $\frac{1}{2} \times 510.7 = 1787$ degrees K. Under the same maximum conditions, the second unit of thermal motion, from T_1 to the solid end point, adds an equal magnitude. Thus the temperature of this theoretical full-scale solid end point is 3575 degrees K. The total expansion coefficient at T_1 on the first segment of the expansion curve, and at the initial point of the second segment, is then $0.0208/3575$. However, this coefficient is subject to a $\frac{1}{9}$ initial level. This makes the net effective coefficient $\frac{8}{9} \times 0.0208/3575 = 5.17 \times 10^{-6}$ per degree K.

Where the end point temperature (which we are equating to the melting point, T_m , for present purposes) is below 3575, the average coefficient of expansion is increased by the ratio $3575/T_m$, inasmuch as the total expansion up to the solid end point is a fixed magnitude. If the first temperature unit, up to T_1 , were to take its full share of the expansion, the coefficient at T_1 on the first segment of the expansion curve, and at the initial point of the second segment, would also be increased by the same ratio. But in the first unit range of temperature the thermal motion takes place in one time region dimension only, and there is no opportunity to increase the total expansion by extension into *additional dimensions* in the manner that is possible when a second unit of motion is involved. (Additional dimensions do not increase the effective magnitude of one unit, as $1^n = 1$.) The total expansion corresponding to the first unit of motion (speed) can be increased by extension to *additional rotational speed displacements*,

but this is possible only in full units, and is limited to a total of four, the maximum magnetic displacement.

As an example, let us consider the element zirconium, which has a melting point of 2125 K. The melting point ratio is $3575/2125 = 1.68$. Inasmuch as this is less than two full units, the expansion coefficient of zirconium remains at one unit (5.17×10^{-6}) at the initial point of the second segment of the curve, and the difference has to be made up by an increase in the rate of expansion between this initial point and T_m ; that is, by an increase in the slope of the second section of the expansion curve. The expansion pattern of zirconium is shown graphically in Fig. 14.

FIGURE 14: *THERMAL EXPANSION*



Now let us look at an element with a lower melting point. Titanium has a melting point of 1941 K. The ratio $3575/1941$ is 1.84. This, again, is less than two full units. Titanium therefore has the same one unit expansion coefficient at the initial level as the elements with higher melting points. The melting point of palladium is only a little more than

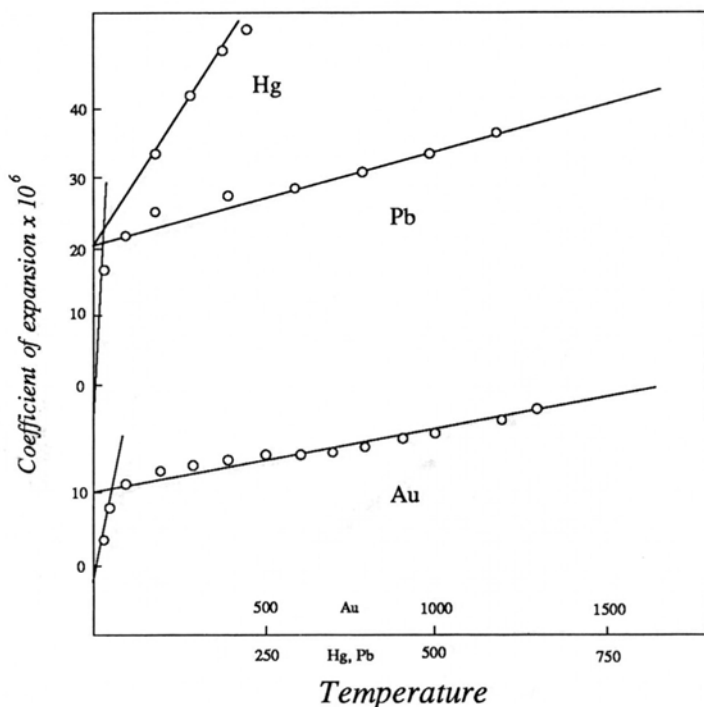
100 degrees below that of titanium, but the difference is just enough to put this element into the two unit range. The ratio computed from the observed melting point 1825 K is 1.96, and is thus slightly under the two unit level. But in this case the difference between the melting point and the end point of the solid state, which we are disregarding for general application, becomes important, as it is enough to raise the 1.96 ratio above 2.00. The expansion coefficient of palladium at the initial point of the second segment of the curve is therefore two units (10.3×10^{-6}), and the expansion follows the pattern illustrated in the second curve in Fig. 14.

The effect of the difference between the solid end point and the melting point can also be seen at the three unit level, as the melting point ratio of silver, $^{3575}/_{1234}=2.90$, is raised enough by this difference to put it over 3.00. Silver then has the three unit (15.5×10^{-6}) expansion coefficient at the upper initial point, as shown in the upper curve in Fig. 14. At the next unit level the element magnesium, with a ratio of 3.87, is similarly near the 4.00 mark, but in this instance the end point increment is not sufficient to close the gap, and magnesium stays on the three unit basis.

None of the elements for which sufficient data were available for comparison with the theoretical curves has a melting point in the four unit range from 715 to 894 K. But since the magnetic rotation is limited to four units, the four unit initial level also applies to the elements with melting points below 715 K. This is illustrated in Fig. 15 by the curve for lead, melting point 601 K.

As can be seen in Fig. 14, the expansion coefficient of silver, as measured experimentally, deviates from the straight line relation in the vicinity of T_1 . This deviation is not due to experimental error or to structural readjustments. It is a result of the nature of the transition from the one unit expansion below T_1 to the multi-unit expansion above this temperature. Unlike the specific heat transition, where the increments represented by the second segment of the specific heat curve *add to* the specific heat at T_1 , the expansion represented by the second segment of the expansion curve *replaces* the expansion represented by the first segment. The initial level of the second segment at zero temperature is the unit (or n-unit) level reached at the end of the first segment.

This means that at T_1 the molecule undergoes an isothermal expansion to the level of the second segment at that temperature. In the aggregate the individual molecular expansions are spread out over a temperature range by the distribution of molecular velocities, and they appear as a bulge in the expansion curve. Coincidentally, there is a downward deviation in the curve, similar to that in the experimental specific heat curves, due to the effect of the transition to the more nearly horizontal second segment of the curve. The net effect of these two types of deviation from the theoretical curve applying to the single molecule depends on their relative magnitude, and on the temperature range over which the deviations are distributed. The curves of Fig. 14 have been selected from among those in which the net deviation is at a minimum, in order to minimize uncertainties in the definition of the upper sections of the curves, and to make it clear that these linear sections actually terminate at the calculated initial levels. More commonly, the bulge is quite pronounced, as in the curves for gold and lead, Fig. 15.

Figure 15: *THERMAL EXPANSION*

When the effect of this systematic deviation from the linear relation in the vicinity of the transition point is taken into consideration, all of the electropositive elements included in the compilation of expansion data utilized in the investigation,¹² except the rare earth elements, have expansion curves that follow the theoretical pattern within the range of accuracy of the experimental results. Most of the rare earths have the one unit expansion coefficient (5.2×10^{-6}) at the initial level of the second segment of the curve, although their melting points are in the range where coefficients of two, or in some cases three, units would be normal. The reason for this, the only deviation from the general pattern in the expansion curves of these elements, is as yet unknown, but it is no doubt connected with the other peculiarities of the rare earth elements that were noted earlier.

The electronegative elements of Division III follow the regular pattern. The lowest melting point in this group is that of mercury, 234K, well below the lowest value for any of the electropositive elements investigated, but this descent to a lower melting point does not introduce any new behavior. The upper segment of the expansion curve for mercury, defined by the empirical data in Fig. 15, definitely terminates at the four unit level (20.7×10^{-6}), as required by the theory. Thus the theoretical relations are applicable to the full temperature range of the first three divisions.

As noted earlier, the borderline elements of Division IV, those with negative electric displacement 4, are capable of acting as members of either Division III or Division IV. The expansion curve for lead, Fig. 15, follows the normal Division III pattern. The lower borderline elements, tin and germanium, have curves in which the initial levels, like those of the rare earths, are lower than the values corresponding to the melting points. Otherwise, these curves are also normal. Very little is known about the expansion of the elements of negative displacement below 4. The theoretical development has not yet been extended to a consideration of the effect of the strongly electronegative character of these elements on the volume relations, and the empirical data are both meager and conflicting.

This Division IV situation is part of the general problem of anisotropic expansion, a subject to which the Reciprocal System

of theory has not yet been applied. The measurements previously cited that apply to anisotropic crystals were made on polycrystalline material in which the expansion in different directions is averaged as a result of the random orientation in the aggregate. Both this issue of anisotropic expansion and the application of the thermal expansion theory to compounds and alloys are still on the waiting list for future investigation. There is no reason to believe that such an investigation will encounter any serious difficulties, but for the present other matters are being given the priority.

CHAPTER 9

Electric Currents

Another set of properties of matter that we will want to consider results from the interaction between matter and one of the sub-atomic particles, the electron. As pointed out in Volume I, the electron, M0-0-(1), in the notation used in this work, is a unique particle. It is the only particle constructed on the material rotational base, M0-0-0, (negative vibration and positive rotation) that has an effective *negative* rotational displacement. More than one unit of negative rotation would exceed the one positive rotational unit of the rotational base, and would result in a negative value of the total rotation. Such a rotation of the basic negative vibration would be unstable in the material environment, for reasons that were explained in the previous discussion. But in the electron the net total rotation is positive, even though it involves one positive and one negative unit, as the positive unit is two-dimensional while the negative unit is one-dimensional.

Furthermore, the independent one-dimensional nature of the rotation of the electron and its positive counterpart, the positron, leads to another unique effect. As we found in our analysis of the rotations that are possible for the basic vibrating unit, the primary rotation of atoms and particles is two-dimensional. The simplest primary rotation has a one-unit magnetic (two-dimensional) displacement, a unit deviation from unit speed, the condition of rest in the physical universe. The electric (one-dimensional) rotation, we found, is not a primary rotation, but merely one that modifies a previously existing two-dimensional rotation. Addition of the one-unit space displacement of the electron rotation to an existing *effective* two-dimensional rotation increases the total scalar speed of rotation. But the one-dimensional rotation of the independent electron does not modify an effective speed; it modifies unit speed, which is zero from

the effective standpoint. The speed displacement of the independent electron, its only effective component, therefore modifies only the effective *space*, not the *speed*.

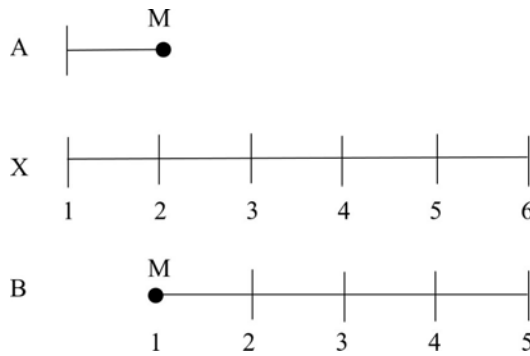
Thus the electron is essentially nothing more than a rotating unit of space. This is a concept that is rather difficult for most of us when it is first encountered, because it conflicts with the idea of the nature of space that we have gained from a long-continued, but uncritical, examination of our surroundings. However, the history of science is full of instances where it has been found necessary to recognize that a familiar, and apparently unique, phenomenon is merely one member of a general class, all members of which have the same physical significance. Energy is a good example. To the investigators who were laying the foundation of modern science in the Middle Ages the property that moving bodies possess by reason of their motion—"impetus" to those investigators; "kinetic energy" to us—was something of a unique nature. The idea that a motionless stick of wood contained the equivalent of this "impetus" because of its chemical composition was as foreign to them as the concept of a rotating unit of space is to most individuals today. But the discovery that kinetic energy is only one form of energy in general opened the door to a major advance in physical understanding. Similarly, the finding that the "space" of our ordinary experience, extension space, as we are calling it in this work, is merely one manifestation of space in general opens the door to an understanding of many aspects of the physical universe, including the phenomena connected with the movement of electrons in matter.

In the universe of motion, the universe whose details we are developing in this work, and whose identity with the observed physical universe we are demonstrating as we go along, space enters into physical phenomena only as a component of motion, and the specific nature of that space is, for most purposes, irrelevant, just as the particular kind of energy that enters into a physical process usually has no relevance to the outcome of the process. The status of the electron as a rotating unit of space therefore gives it a very special role in the physical activity of the universe. It should be noted at this time that the electron that we are now discussing carries no charge. It is a combination of two motions, a basic vibration and a rotation

of the vibrating unit. As we will see later, an electric charge is an additional motion that *may* be superimposed on this two-component combination. The behavior of charged electrons will be considered after some further groundwork has been laid. For the present we are concerned only with the uncharged electrons.

As a unit of space, the uncharged electron cannot move through extension space, since the relation of space to space does not constitute motion. But under appropriate conditions it can move through ordinary matter, inasmuch as this matter is a combination of motions with a net positive, or time, displacement, and the relation of space to time does constitute motion. The present-day view of the motion of electrons in solid matter is that they move through the spaces between the atoms. The resistance to the electron flow is then considered to be analogous to friction. Our finding is that the electrons (units of space) exist *in* the matter, and move *through* that matter in the same manner as that matter moves through extension space.

The motion of the electrons is negative, or outward, with respect to the net motion of material objects. This is illustrated in the following diagram:



Line X in the diagram is a representation of a scalar magnitude of extension space, as it appears in the conventional reference system. Line A shows the effect of a unit of motion of a material object M through that space. The object that was originally coincident with spatial unit 1 is now coincident with spatial unit 2. Line B shows what happens if the original motion of object M is followed by a unit of electron motion through M. Just as object M moved through space X in line A, so space X (the electrons) moves through object

M in line B. In one unit of motion (line A) object M advances from spatial unit 1 to spatial unit 2. In the following unit of the inverse type of motion (line B) the numbered spatial locations advance one unit relative to object M. This brings M back into coincidence with spatial unit 1, the same result that would have followed if object M had moved backward in the absence of any electron movement. Thus the movement of space (electrons) through matter is equivalent to a negative movement of matter through space. It follows that the voltage differential that causes the electron motion, and the stress in any substance that absorbs the motion, are likewise negative.

Directional movement of electrons through matter will be identified as an *electric current*. If the atoms of the matter through which the current passes are effectively at rest relative to the structure of the solid aggregate as a whole, uniform motion of the electrons (space) through matter has the same general properties as motion of matter through space. It follows Newton's first law of motion, and can continue indefinitely without addition of energy. This situation exists in the phenomenon known as *superconductivity* that has been observed experimentally in many substances at very low temperatures. But where the atoms of a material aggregate are in effective motion thermally, movement of electrons through the matter adds to the spatial component of the thermal motion (that is, increases the speed), and thereby imparts energy (heat) to the moving atoms.

The magnitude of the current is measured by the number of electrons (units of space) per unit of time. Units of space per unit of time is the definition of speed, hence the electric current is a speed. From a mathematical standpoint it is immaterial whether a mass is moving through extension space or space is moving through the mass. Thus in dealing with the electric current we are dealing with the *mechanical* aspects of electricity, and the current phenomena can be described by the same mathematical equations that are applicable to ordinary motion in space, with appropriate modifications for differences in conditions, where such differences exist. It would even be possible to use the same units, but for historical reasons, and as a matter of convenience, a separate system of units is utilized in present-day practice.

The basic unit of current electricity is the unit of quantity. In the natural system it is the spatial aspect of one electron, which has a

speed displacement of one unit. Quantity, q , is therefore equivalent to space, s . Energy has the same status in current flow as in the mechanical relations, and has the space-time dimensions $\frac{1}{2}$. Energy divided by time is power, $\frac{1}{2}$. A further division by current, which has the dimensions of speed, $\frac{1}{4}$, then produces electromotive force (emf) with the dimensions $\frac{1}{2} \times \frac{1}{2} = \frac{1}{2}$. These are, of course, the space-time dimensions of force in general.

The term “electric potential” is commonly used as an alternative to emf, but, for reasons to be discussed later, we will not use “potential” in this sense. Where a more convenient term than emf is appropriate, we will use the term “voltage,” symbol V .

Dividing voltage, $\frac{1}{2}$, by current, $\frac{1}{4}$, we obtain $\frac{1}{2} \times \frac{4}{1} = 2$. This is resistance, symbol R , the only electrical quantity thus far considered that is not equivalent to a familiar mechanical quantity. Its true nature is revealed by an examination of its space-time structure. The dimensions $\frac{1}{2}$ are equivalent to mass, $\frac{1}{2}$, divided by time, t . Resistance is therefore mass per unit time. The relevance of such a quantity can easily be seen when it is realized that the amount of mass entering into the motion of space (electrons) through matter is not a fixed quantity, as it is in the motion of matter through extension space, but a quantity whose magnitude depends on the amount of movement of the electrons. In motion of matter through extension space the mass is constant while the space depends on the duration of the movement. In the current flow the space (number of electrons) is constant while the mass depends on the duration of the movement. If the flow is only momentary, each electron may move through only a small fraction of the total amount of mass in the circuit, whereas if it continues for a longer time the entire circuit may be traversed repeatedly. In either case, the total *mass* involved in the current flow is the product of the mass per unit time (the resistance) and the time of flow. In the movement of matter through extension space, the total *space* is determined in the same manner; that is, it is the product of the space per unit time (velocity) and the time of movement.

In dealing with resistance as a property of matter we will be interested mainly in the *specific resistance*, or *resistivity*, which is defined as the resistance of a unit cube of the substance under consideration. Resistance is directly proportional to the distance

traveled by the current and inversely proportional to the cross-sectional area of the conductor. It follows that if we multiply the resistance by unit area and divide by unit distance we arrive at a quantity with the dimensions t^2/s^2 that reflects only the inherent characteristics of the material and the environmental conditions (principally temperature and pressure) and is independent of the geometrical structure of the conductor. The reciprocals of resistance and resistivity are *conductance* and *conductivity*, respectively.

With the benefit of the clarification of the space-time dimensions of resistance we can now go back to the empirically determined relations between resistance and other electrical quantities, and verify the consistency of the space-time identifications.

Voltage:	$V = IR = \text{s/t} \times \text{t}^2/\text{s}^3 = \text{t}/\text{s}^2$
Power:	$P = I^2R = \text{s}^2/\text{t}^2 \times \text{t}^2/\text{s}^3 = 1/\text{s}$
Energy:	$E = I^2Rt = \text{s}^2/\text{t}^2 \times \text{t}^2/\text{s}^3 \times \text{t} = \text{t}/\text{s}$

The energy equation demonstrates the equivalence of the mathematical expressions of the electrical and mechanical phenomena. Since resistance is mass per unit time, the product of resistance and time, Rt , is equivalent to mass, m . The current, I , is a speed, v . The electrical energy expression RtI^2 is thus dimensionally equivalent to the kinetic energy expression $\frac{1}{2}mv^2$. In other words, the quantity RtI^2 is the kinetic energy of the electron motion.

Instead of using resistance, time, and current, we may put the energy expression into terms of voltage, V (equivalent to IR), and quantity, q , (equivalent to It). The expression for the magnitude of the energy (or work) is then $W = Vq$. Here we have a definite confirmation of the identification of electric quantity as the equivalent of space. Force, as described in one of the standard physics textbooks, is “an explicitly definable vector quantity that changes or tends to produce a change in the motion of bodies.” Electromotive force, or voltage, conforms to this description. It tends to cause motion of the electrons in the direction of the voltage gradient. Energy in general is the product of force and distance. Electrical energy, as Vq , is the product of force and quantity. It follows that electrical quantity is equivalent to distance: the same conclusion that we derived from the theoretical nature of the uncharged electron.

In conventional scientific thought the status of electrical energy as one form of energy in general is accepted, as it must be, since it can be converted to any of the other forms, but the status of electrical, or electromotive, force as one form of force in general is *not* accepted. If it were, the conclusion stated in the preceding paragraph would be inescapable. But the clear verdict of the observed facts is disregarded because there is a general impression that electrical quantity and space are entities of a totally different nature.

The early investigators of electrical phenomena recognized that the quantity measured in volts has the characteristics of a force, and they named it accordingly. Contemporary theorists reject this identification because it conflicts with their views as to the nature of the electric current. W. J. Duffin, for instance, gives us a definition of electromotive force (emf), and then says,

In spite of its name, it is clearly not a force but is equal to the work done per unit positive charge in taking a charge completely around [the electric circuit]; its unit is therefore the volt.¹³

Work per unit of space is force. This author simply takes it for granted that the moving entity, which he calls a charge, is *not* equivalent to space, and he therefore deduces that the quantity measured in volts cannot be a force. Our finding is that his assumptions are wrong, that the moving entity is not a charge, but is a rotating unit of space (an uncharged electron). The electromotive force, measured in volts, is then, in fact, a force. In effect, Duffin concedes this point when he tells us, in another connection, that “V/m [volts per meter] is the same as N/C [newtons per coulomb].”¹⁴ Both express the voltage gradient in terms of force divided by space.

Conventional physical theory does not pretend to give us any understanding of the nature of either electrical quantity or electric charge. It simply *assumes* that inasmuch as scientific investigation has hitherto been unable to produce any explanation of its nature, the electric charge must be a unique entity, independent of the other fundamental physical entities, and must be accepted as one of the “given” features of nature. It is then further *assumed* that this entity of unknown nature that plays the central role in electrostatic phenomena is identical with the entity of unknown nature, electrical quantity, that plays the central role in current electricity.

The most significant weakness of the conventional theory of the electric current, the theory based on the foregoing assumptions, as we now see it in the light of the more complete understanding of physical fundamentals derived from the theory of the universe of motion, is that it assigns two different, and incompatible, roles to the electrons. These particles, according to present-day theory, are *components* of the atomic structure, yet at least some of them are presumed to be free to accommodate themselves to any electrical forces applied to the conductor. On the one hand, each is so firmly bound to the remainder of the atom that it plays a significant part in determining the properties of that atom, and a substantial force (the ionization potential) must be applied in order to separate it from the atom. On the other hand, these electrons are so free to move that they will respond to thermal or electrical forces whose magnitude is only slightly above zero. They must exist in a conductor in specific numbers in order to account for the fact that the conductor is electrically neutral while carrying current, but at the same time they must be free to leave the conductor, either in large or small quantities, if they acquire sufficient kinetic energy.

It should be evident that the theories are calling upon the electrons to perform two different and contradictory functions. They have been assigned the key position in both the theory of atomic structure and the theory of the electric current, without regard for the fact that the properties that they must have in order to perform the functions required by either one of these theories disqualify them for the functions that they are called upon to perform in the other.

In the theory of the universe of motion, each of these phenomena involves a different physical entity. The unit of atomic structure is a unit of rotational motion, not an electron. It has the quasi-permanent status that is required of an atomic constituent. The electron, without a charge, and without any connection with the atomic structure, is then available as the freely moving unit of the electric current.

The fundamental postulate of the Reciprocal System of theory is that the physical universe is a universe of motion, one in which all entities and phenomena are motions, combinations of motions, or relations between motions. In such a universe none of the basic phenomena are unexplainable. "Unanalyzables," as Bridgman called them, do

not exist. The basic physical entities and phenomena of the universe of motion—radiation, gravitation, matter, electricity, magnetism, and so on—can be defined explicitly in terms of space and time. Unlike conventional physical theory, the Reciprocal System does not have to leave its basic elements cloaked in metaphysical mystery. It does not have to exclude them from physical inquiry, in the manner of the following statement from the *Encyclopedia Britannica*:

The question “What is electricity?” like the question “What is matter?” really lies outside the realm of physics and belongs to that of metaphysics.¹⁵

In a universe composed entirely of motion, an electric charge applied to a physical entity must necessarily be a motion. Thus the problem faced in the theoretical investigation was not to answer the question, *What* is an electric charge?, but merely to determine *what kind of motion* manifests itself as a charge. The identification of the charge as an added motion not only clarifies the relation between the charged electron that is observed experimentally and the uncharged electron that is known only as the moving entity in the electric current, but also explains the interchanges between the two that are the principal support for the currently popular opinion that only one entity, the charge, is involved. It is not always remembered that this opinion achieved general acceptance only after a long and spirited controversy. There are similarities between static and current phenomena, but there are also significant differences. Inasmuch as no theoretical explanation of either kind of electric effect was available at the time, the question to be decided was whether to regard the two as identical because of the similarities, or as disparate because of the differences. Once made, the decision in favor of identity has persisted, even though much evidence against its validity has accumulated in the meantime.

The similarities are of two general types: (1) some of the properties of charged particles and electric currents are alike, and (2) there are observable transitions from one to the other. Identification of the charged electron as an uncharged electron with an added motion explains both types of similarities. For instance, a demonstration that a rapidly moving charge has the same magnetic properties as an electric current was a major factor in the victory won by the proponents of the “charge” theory of the electric current many years

ago. But our findings are that the moving entities are electrons, or other *carriers* of the charges, and the existence or non-existence of electric charges is irrelevant. The second kind of evidence that has been interpreted as supporting the identity of the static and current electrons is the apparent continuity from the electron of current flow to the charged electron in such processes as electrolysis. Here the explanation is that the electric charge is easily created and easily destroyed. As everyone knows, nothing more than a slight amount of friction is sufficient to generate an electric charge on many surfaces, such as those of present-day synthetic fabrics. It follows that wherever a concentration of energy exists in one of these forms that can be relieved by conversion to the other, the rotational vibration that constitutes a charge is either initiated or terminated in order to permit the type of electron motion that can take place in response to the effective force.

It has been possible to follow the prevailing policy, regarding the two different quantities as identical, and utilizing the same units for both, only because the two different usages are entirely separate in most cases. Under these circumstances no error is introduced into the calculations by using the same units, but a clear distinction is necessary in any case where either calculation or theoretical consideration involves quantities of both kinds.

As an analogy we might assume that we are undertaking to set up a system of units in which to express the properties of water. Let us further assume that we fail to recognize that there is a difference between the properties of weight and volume, and consequently express both in cubic centimeters. Such a system is equivalent to using a weight unit of one gram, and as long as we deal separately with weight and volume, each in its own context, the fact that the expression "cubic centimeter" has two entirely different meanings will not result in any difficulties. However, if we have occasion to deal with both quantities simultaneously, it is essential to recognize the difference between the two. Dividing cubic centimeters (weight) by cubic centimeters (volume) does not result in a pure number, as the calculations seem to indicate; the quotient is a *physical* quantity with the dimensions weight/volume. Similarly, we may use the same units for electric charge and electric quantity as long as they are

employed independently and in the right context, but whenever the two enter into the same calculation, or are employed individually with the wrong physical dimensions, there is confusion.

This dimensional confusion resulting from the lack of distinction between the charged and uncharged electrons has been a source of considerable concern, and some embarrassment to the theoretical physicists. One of its major effects has been to prevent setting up any comprehensive systematic relationship between the dimensions of physical quantities. The failure to find a basis for such a relationship is a clear indication that something is wrong in the dimensional assignments, but instead of recognizing this fact, the current reaction is to sweep the problem under the rug by pretending that it does not exist. As one observer sees the picture:

In the past the subject of dimensions has been quite controversial. For years unsuccessful attempts were made to find “ultimate rational quantities” in terms of which to express all dimensional formulas. It is now universally agreed that there is no one “absolute” set of dimensional formulas.¹⁶

This is a very common reaction to long years of frustration, one that we encountered frequently in our examination of the subjects treated in Volume I. When the most strenuous efforts by generation after generation of investigators fail to reach a defined objective, there is always a strong temptation to take the stand that the objective is inherently unattainable. “In short,” says Alfred Landé, “if you cannot clarify a problematic situation, declare it to be ‘fundamental,’ then proclaim a corresponding ‘principle’.”¹⁷ So physical science fills up with principles of impotence rather than explanations.

In the universe of motion the dimensions of *all* quantities of *all* kinds can be expressed in terms of space and time only. The space-time dimensions of the basic mechanical quantities were identified in Volume I. In this chapter we have added those of the quantities involved in the flow of electric current. The chapters that follow will complete this task by identifying the space-time dimensions of the magnetic quantities and the additional electric quantities that make their appearance in the phenomena due to charges of one kind or another and in the magnetic effects of electric currents.

This clarification of the dimensional relations is accompanied by a determination of the natural unit magnitudes of the various physical quantities. The system of units commonly utilized in dealing with electric currents was developed independently of the mechanical units on an arbitrary basis. In order to ascertain the relation between this arbitrary system and the natural system of units it is necessary to measure some one physical quantity whose magnitude can be identified in the natural system, as was done in the previous determination of the relations between the natural and conventional units of space, time, and mass. For this purpose we will use the *Faraday constant*, the observed relation between the quantity of electricity and the mass involved in electrolytic action. Multiplying this constant, 2.89366×10^{14} esu/g-equiv., by the natural unit of atomic weight, 1.65979×10^{-24} g, we arrive at 4.80287×10^{-10} esu as the natural unit of electrical quantity.

The magnitude of the electric current is the number of electrons per unit of time; that is, units of space per unit of time, or speed. Thus the natural unit of current could be expressed as the natural unit of speed, 2.99793×10^{10} cm/sec. In electrical terms it is the natural unit of quantity divided by the natural unit of time, 1.520655×10^{-16} sec, and amounts to 3.15842×10^6 esu/sec, or 1.05353×10^{-3} amperes. The conventional unit of electrical energy, the watt-hour, is equal to 3.6×10^{10} ergs. The natural unit of energy, 1.49275×10^{-3} ergs, is therefore equivalent to 4.14375×10^{-14} watt-hours. Dividing this unit by the natural unit of time, we obtain the natural unit of power 9.8099×10^{12} ergs/sec = 9.8099×10^5 watts. A division by the natural unit of current then gives us the natural unit of electromotive force, or voltage, 9.31146×10^8 volts. Another division by current brings us to the natural unit of resistance, 8.83834×10^{11} ohms.

The basic quantities of current electricity and their natural units in electrical terms can be summarized as follows:

s	quantity	4.80287×10^{-10} esu
$\frac{s}{t}$	current	1.05353×10^{-3} amperes
$\frac{1}{s}$	power	9.8099×10^5 watts
$\frac{t}{s}$	energy	4.14375×10^{-14} watt-hours
$\frac{t}{s^2}$	voltage	9.31146×10^8 volts
$\frac{t^2}{s^3}$	resistance	8.83834×10^{11} ohms

Another electrical quantity that should be mentioned because of the key role that it plays in the present-day mathematical treatment of magnetism is “current density,” which is defined as “the quantity of charge passing per second through unit area of a plane normal to the line of flow.” This is a strange quantity, quite unlike any physical quantity that has previously been discussed in the pages of this and the preceding volume, in that it is not a relation *between* space and time. When we recognize that the quantity is actually *current* per unit of area, rather than “charge” (a fact that is verified by the units, amperes per square meter, in which it is expressed), its space-time dimensions are seen to be $s_t \times 1/s^2 = 1/s_t$. These are not the dimensions of a motion, or a property of a motion. It follows that this quantity, as a whole, has no physical significance. It is merely a mathematical convenience.

The fundamental laws of current electricity known to present-day science—such as Ohm’s Law, Kirchhoff’s Laws, and their derivatives—are empirical generalizations, and their application is not affected by the clarification of the essential nature of the electric current. The substance of these laws, and the relevant details, are adequately covered in existing scientific and technical literature. In conformity with the general plan of this work, as set forth earlier, these subjects will not be included in our presentation.

This is an appropriate time to make some comments about the concept of “natural units.” There is no ambiguity in this concept, so far as the basic units of motion are concerned. The same is true, in general, of the units of the simple scalar quantities, although some questions do arise. For example, the unit of space in the region inside unit distance, the time region, as we are calling it, is inherently just as large as the unit of space in the region outside unit distance, but *as measured* it is reduced by the inter-regional ratio, 156.444, for reasons previously explained. We cannot legitimately regard this quantity as something less than a full unit, since, as we saw in Volume I, particularly in connection with the structure of organic compounds, it has the same status in the time region that the full-sized natural unit of space has in the region outside unit distance. On the other hand, we have to recognize that this unit enters into the quantitative aspects of most physical relations as a smaller magnitude. The logical way of handling this situation appears to be

to take the stand that there are two different natural units of distance (one-dimensional space), a simple unit and a compound unit, that apply under different circumstances.

The more complex physical quantities are subject to still more variability in the unit magnitudes, because these quantities are combinations of the simpler quantities, and the combination may take place in different ways and under different conditions. For instance, as we saw in our examination of the units of mass in Volume I, there are several different manifestations of mass, each of which involves a different combination of natural units and therefore has a natural unit of its own. In this case, the primary cause of variability is the existence of a secondary mass component that is related to the primary mass by the inter-regional ratio, or a modification thereof, but some additional factors introduce further variability, as indicated in the earlier discussion.

CHAPTER 10

Electrical Resistance

While the motion of the electric current through matter is equivalent to motion of matter through space, as brought out in the discussion in Chapter 9, the conditions under which each type of motion is encountered in our ordinary experience emphasize different aspects of the common fundamentals. In dealing with the motion of matter through extension space we are primarily concerned with the motions of individual objects. Newton's laws of motion, which are the foundation stones of mechanics, deal with the application of forces to initiate or modify the motions of such objects, and with the transfer of motion from one object to another. Our interest in the electric current, on the other hand, is concerned mainly with the *continuous* aspects of the current flow, and the status of the individual objects that are involved is largely irrelevant.

The mobility of the spatial units in the current flow also introduces some kinds of variability that are not present in the movement of matter through extension space. Consequently, there are behavior characteristics, or properties, of material structures that are peculiar to the relation between these structures and the moving electrons. Expressing this in another way, we may say that matter has some distinctive *electrical properties*. The basic property of this nature is resistance. As pointed out in Chapter 9, resistance is the only quantity participating in the fundamental relations of current flow that is not a familiar feature of the mechanical system of equations, the equations that deal with the motion of matter through extension space.

Present-day ideas as to the origin of electrical resistance are summarized by one author in this manner:

Ability to conduct electricity... is due to the presence of large numbers of quasi-free electrons which under the action of an

applied electric field are able to flow through the metallic lattice...
Disturbing influences...impede the free flow of electrons,
scattering them and giving rise to a resistance.¹⁸

As indicated in the preceding chapter, the development of the theory of the universe of motion arrives at a totally different concept of the nature of electrical resistance. The electrons, we find, are derived from the environment. It was brought out in Volume I that there are physical processes in operation which produce electrons in substantial quantities, and that, although the rotational motions that constitute these electrons are, in many cases, absorbed by atomic structures, the opportunities for utilizing this type of motion in such structures are limited. It follows that there is always a large excess of free electrons in the material sector of the universe, most of which are uncharged. In this uncharged state the electrons cannot move with respect to extension space, because they are inherently rotating units of space, and the relation of space to space is not motion. In open space, therefore, each uncharged electron remains permanently in the same location with respect to the natural reference system, in the manner of a photon. In the context of the stationary spatial reference system the uncharged electron, like the photon, is carried outward at the speed of light by the progression of the natural reference system. All material aggregates are thus exposed to a flux of electrons similar to the continual bombardment by photons of radiation. Meanwhile there are other processes, to be discussed later, whereby electrons are returned to the environment. The electron population of a material aggregate such as the earth therefore stabilizes at an equilibrium level.

These processes that determine the equilibrium electron concentration are independent of the nature of the atoms of matter and of the atomic volume. The concentration of electrons is therefore uniform in electrically isolated conductors where there is no current flow. It follows that the number of electrons involved in the thermal motion of atoms of matter is proportional to the atomic volume, and the energy of that motion is determined by the effective rotational factors of the atoms. The atomic volume and thermal energy therefore determine the resistance.

Those substances whose rotational motion is entirely in time (Divisions I and II) have their thermal motion in space, in accordance with the general rule governing addition of motions, as set forth in Volume I. For these substances zero thermal motion corresponds to zero resistance, and the resistance increases with the temperature. This is due to the fact that the concentration of electrons (units of space) in the time component of the conductor is constant for any specific current magnitude. The current therefore increases the thermal motion by a specific proportion. Such substances are *conductors*.

Where there are two dimensions of rotation in space, as in many of the elements of Division IV, the thermal motion, which requires two open dimensions because of the finite diameter of the moving electron, is necessarily in time. In this case, zero temperature corresponds to zero motion in time. Here the resistance is initially extremely high, but decreases with an increase in temperature. Substances of this kind are known as *insulators* or *dielectrics*.

Where there is only one dimension of spatial rotation, as in Division III, the elements of greatest electric displacement, those closest to the electropositive divisions, are able to follow the positive pattern, and are conductors. The Division III elements of lower electric displacement follow a modified time motion pattern, with resistance decreasing from a high, but finite, level at zero temperature. These substances of intermediate characteristics are *semiconductors*.

For the present we will be concerned primarily with the resistance of conductors, and will further limit the discussion to what may be called the “regular” pattern of conductor resistance. A limitation of this kind is necessary at the present stage of the investigation because the large element of uncertainty in the experimental information on the resistivity of the various conducting materials makes the clarification of the resistance relations a slow and difficult process. The early stages of the development of the Reciprocal System of theory, prior to the publication of the first edition of this work in 1959, which were very productive in the non-electrical areas, made relatively little progress in dealing with the electrical properties, largely because of conflicts between the theoretical deductions and

some experimental results that have since been found to be incorrect. The increasing scope and accuracy of the experimental work in the intervening years has improved this situation very materially, but the basic problem still remains.

Ideally it should be possible to deduce all of the pertinent information from theoretical premises alone, without reference to experimental determinations, but as a practical matter this is not feasible. A few steps can be, and have been, taken on a purely theoretical basis, particularly where the previous development of the theory has cast some important new light on the subject matter, but from the practical standpoint an extensive and detailed investigation in any area is possible only if the theoretical study and the checking of the theoretical conclusions against experimental and observational data go hand in hand. It follows that where empirical data are lacking, progress is difficult, and where they are seriously wrong, it is essentially impossible.

Unfortunately, resistance measurements are subject to many factors that introduce uncertainty into the results. The purity of the specimen is particularly critical because of the great difference between the resistivities of conductors and dielectrics. Even a very small amount of a dielectric impurity can alter the resistance substantially. Conventional theory has no explanation for the magnitude of this effect. If the electrons move through the interstices between the atoms, as this theory contends, a few additional obstacles in the path should not contribute significantly to the resistance. But, as we saw in Chapter 9, the current moves *through* all of the atoms of the conductor, including the impurity atoms, and it increases the heat content of each atom in proportion to its resistance. The extremely high dielectric resistance results in a large contribution by each impurity atom, and even a very small number of such atoms therefore has a significant effect. Semiconducting elements are less effective as impurities, because of their lower resistance, but they may still have resistivities thousands of times as great as those of the conductor metals.

The resistance also varies under heat treatment, and careful annealing is required before reliable measurements can be made. The adequacy of this treatment in many, if not most, of the resistance determinations is questionable. For example, G. T. Meaden reports that the resistance of beryllium was lowered more than fifty percent

by such treatment, and comments that “much earlier work was clearly based on unannealed specimens.”¹⁹ Other sources of uncertainty include changes in crystal structure or magnetic behavior that take place at different temperatures or pressures in different specimens, or under different conditions, often with substantial hysteresis effects.

Ultimately, of course, it will be desirable to bring all of these variables within the scope of the theoretical treatment, but for the present our objective will have to be limited to deducing from the theory the nature and magnitude of the changes in resistance resulting from temperature and pressure variations in the absence of these complicating factors, and then to show that enough of the experimental results are in agreement with the theory to make it probable that the discrepancies, where they occur, are due to one of more of these factors that modify the normal values.

Inasmuch as the electrical resistance is a product of the thermal motion, the energy of the electron motion is in equilibrium with the thermal energy. The resistance is therefore directly proportional to the effective thermal energy; that is, to the temperature. It follows that the increment of resistance per degree is a constant for each (unmodified) substance, a magnitude that is determined by the atomic characteristics. The curve representing the relation of the resistivity to the temperature, in application to a single atom, is thus linear. Like the curves representing the temperature variation of the other properties that we examined in the earlier chapters, and for the same reasons, the initial level of the resistivity curve is negative. From this initial level to the melting point the resistivity of an unmodified atom (one that has not undergone a structural rearrangement or other change that modifies the resistance relations) follows a single straight line, rather than a curve composed of two or more segments of different slopes, as in the specific heat and thermal expansion curves. This limitation to a single line is characteristic of the electron relations, and is due to the fact that the electron has only one rotational displacement unit, and therefore cannot shift to a multi-unit type of motion in the manner of the complex atomic structures.

A somewhat similar change in the resistivity curve does occur, however, if the factors that determine the resistance are modified by

some rearrangement of the kind mentioned earlier. As P. W. Bridgman commented in discussing some of his results, after a change of this nature has taken place, we are really dealing with a different substance. The curve for the modified atom is also a straight line, but it is not collinear with the curve of the unmodified atom. At the point of transition to the new form the resistivity of the individual atom abruptly changes to a different straight line relation. The resistivity of the aggregate follows a transition curve from one line to the other, as usual. At the lower end of the temperature range, the resistivity of the solid aggregate follows another transition curve of the same nature as those that we found in the curves representing the properties discussed earlier. The relation of the resistance to the temperature in this temperature range is currently regarded as exponential, but as we saw in other cases of the same kind, it is actually a probability curve that reflects the resistivity of the diminishing number of atoms that are still individually above the temperature at which the atomic resistivity reaches the zero level. The curve for the solid aggregate also diverges from the single atom curve at the upper end, due to the increasing proportion of liquid molecules in the solid aggregate.

In this case, again, two values are required for a complete definition of the linear curve; either the coordinates of two points on the curve, or the slope of the curve and the location of one fixed point. A fixed point that is available from theoretical premises is the zero point temperature, the point at which the curve for the individual atom reaches the zero resistance level. The theoretical factors that determine this temperature are the same as those applying to the specific heat and thermal expansion curves, except that since the resistivity is an interaction between the atom and the electron it is effective only when the motions of both objects are directed outward. The theoretical zero point temperature normally applicable to resistivity is therefore twice that applicable to the properties previously considered.

Up to this point the uncertainties in the experimental results have had no effect on the comparison of the theoretical conclusions with experience. It is conceded that the relation of resistivity to temperature is generally linear, with deviations from linearity in certain temperature ranges and under certain conditions. The only question at issue is whether these deviations are satisfactorily

explained by the Reciprocal System of theory. When this question is considered in isolation, without taking into account the status of that system as a *general* physical theory, the answer is a matter of judgment, not a factual matter that can be resolved by comparison with observation. But we have now arrived at a place where the theory identifies some specific numerical values. Here agreement between theory and observation is a matter of objective fact, not one that calls for a judgment. But agreement within an acceptable margin can be expected only if (1) the experimental resistivities are reasonably accurate, (2) the zero point temperatures applicable to specific heat (which are being used as a base) were correctly determined, and (3) the theoretical calculation and the resistivity measurement refer to the same structure.

Table 24 applies equation 7-1, with a doubled numerical constant, and the rotational factors from Table 22, to a determination of the temperatures of the zero levels of the resistance curves of the elements included in the study, and compares the results with the corresponding points on the empirical curves. The amount of uncertainty in the resistivity measurements is reflected in the fact that for 11 of these 40 elements there are two sets of experimental results that have been selected as the “best” values by different data compilers.²⁰ In three other cases there are substantial differences in the experimental results at the higher temperatures, but the curves converge on the same value of the zero resistivity temperature. In a situation where uncertainties of this magnitude are prevalent, it can hardly be expected that there will be anywhere near a complete agreement between the theoretical and experimental values. Nevertheless, if we take the closer of the two “best” experimental results in the 11 two-value cases, the theoretical and experimental values agree within four degrees in 26 of the 40 elements, almost two-thirds of the total.

The rare earth elements were not included in this study because the resistances of these elements, like so many of their other properties, follow a pattern differing in some respects from that of most other elements, including a transition to a new structural form at a relatively low temperature, accompanied by a major decrease in the slope of the resistivity curve. Because of this low temperature transition it is

difficult to locate the zero point temperature from the empirical data, but in 9 of the 13 elements of this group for which sufficient data are available to enable an approximate identification of this temperature, it appears to be between 10 and 20 degrees K. The theoretical range for these elements, as indicated by the factors listed in Table 22, is from 12 to 20 degrees. Here again, then, the measured resistivities of two-thirds of the elements are at least approximately in agreement with the theoretical values.

TABLE 24: *TEMPERATURE OF ZERO RESISTANCE*

	Total Factors	T ₀			Total Factors	T ₀	
		Calc.	Obs.			Calc.	Obs.
Li	14	56	56	Ru	14	56	44-59
Na	6	24	30	Rh	13	52	44-55
Mg	12	48	45	Pd	10	40	39
Al	14	56	57-60	Ag	8	32	28-35
K	4	16	17	Cd	5	20	18
Sc	10	40	33	In	12	48	19
Ti	14	56	54	Sn	7	28	25
V	12	48	45	Sb	8	32	24-35
Cr	14	56	69	Cs	2	8	8
Fe	16	64	73	Ba	4	16	26
Co	14	56	64-78	Hf	8	32	32
Ni	14	56	55	Ta	8	32	30
Cu	12	48	46-49	W	12	48	46-55
Zn	8	32	27	Re	10	40	45
Ga	4	16	31	Ir	11	44	28-46
As	12	48	42	Pt	8	32	33
Rb	2	8	11	Au	6	24	18
Y	8	32	28	Hg	4	16	7
Zr	9	36	30-45	Tl	4	16	16
Mo	14	36	36-55	Pb	4	16	12

The existence of this amount of agreement, in spite of all of the influences tending to generate discrepancies, is about as good a confirmation of the validity of the theory, as a general proposition,

as can be expected under the existing circumstances. Furthermore, it is not unlikely that there are alternate resistance patterns that result in explainable deviations from the calculated values, and some of the larger discrepancies may be thus accounted for when an investigation of broader scope is undertaken.

For the second defining value of the resistivity curves we can use the temperature coefficient of resistivity, the slope of the curve, a magnitude that reflects the inherent resistivity of the conductor material. The temperature coefficient as given in the published physical tables is not the required value. This is merely a relative magnitude, the incremental change in resistivity relative to the resistivity at a reference temperature, usually 20 degrees C. What is needed for present purposes is the absolute coefficient, in microhm-centimeters per degree, or some similar unit.

Some studies have been made in this area, and as might be expected, it has been found that the electric (one-dimensional) speed displacement is the principal determinant of the resistivity, in the sense that it is responsible for the greatest amount of variation. However, the effective quantity is not usually the normal electric displacement of the atoms of the element involved, as this value is generally modified by the way that the atom interacts with the electrons. The conclusions that have been reached as to the nature and magnitude of these modifications are still rather tentative, and there are major uncertainties in the empirical values against which the theoretical results would normally be checked to test their validity. The results of these studies have therefore been omitted from this volume, in conformity with the general policy of restricting the present publication to those results whose validity is firmly established.

The experimental difficulties that introduce uncertainties into the correlations between the theoretical and experimental values of the resistivity do not play as large a role in the relative resistance under compression. The compression results therefore give us a more definite and unequivocal picture. Again, however, this initial exploration of the subject, as it appears in the context of the Reciprocal System of theory, will have to be confined to the "regular" pattern, the one followed by most of the metallic conductors.

Because the movement of electrons (space) through matter is the inverse of the movement of matter through space, the inter-regional relations applicable to the effect of pressure on resistance are the inverse of those that apply to the change in volume under pressure. We found in Chapter 4 that the volume of a solid under compression conforms to the relation $PV^2=k$. By reason of the inverse nature of the electron movement, the corresponding equation for electrical resistance is

$$P^2R=k \quad (10-1)$$

As in the compressibility equation, the symbol P in this expression refers to the total effective pressure. If we give the internal component of this total the designation P_0 , as in the volume compressibility discussion, and limit the term P to the externally applied pressure, the equation becomes

$$(P + P_0)^2 R = k \quad (10-2)$$

The general situation with respect to the values of the internal pressure applicable to resistance is essentially the same as that encountered in the study of compressibility. Some elements maintain the same internal pressure throughout Bridgman's entire pressure range, some undergo second order transitions to higher P_0 values, and others are subject to first order transitions, just as in the volume relations. However, the internal pressure applicable to resistance is not necessarily the same as that applicable to volume. In some substances, tungsten and platinum, for example, these internal pressures actually are identical at every point in the pressure range of Bridgman's experiments. In another, and larger, class, the applicable values of P_0 are the same as in compression, but the transition from the lower to the higher pressure takes place at a different temperature.

The values for nickel and iron illustrate this common pattern. The initial reduction in the volume of nickel took place on the basis of an internal pressure of 913 M kg/cm². Somewhere between an external pressure of 30 M kg/cm² (Bridgman's pressure limit on this element) and 100 M kg/cm² (the initial point of later experiments at very high pressure) the internal pressure increased to 1370 M kg/cm² (from azy factors 4-8-1 to 4-8-1½). In the resistance measurements the same

transition occurred, but it took place at a lower external pressure, between 10 and 20 M kg/cm². Iron has the same internal pressures in resistance as nickel, with the transition at a somewhat higher external pressure, between 40 and 50 kg/cm². But in compression this transition did not appear at all in Bridgman's pressure range, and was evident only in the shock wave experiments carried to much higher pressures.

Table 25 is a comparison of the internal pressures in resistance and compression for the elements included in the study. The symbol x following or preceding some of the values indicates that there is evidence of a transition to or from a different internal pressure, but the available data are not sufficient to define the alternate pressure level.

TABLE 25: *INTERNAL PRESSURES IN RESISTANCE AND COMPRESSION*
(Bridgman's pressure range)

	P ₀ (M kg/cm ²)			P ₀ (M kg/cm ²)	
	Compression	Resistance		Compression	Resistance
Be	571-856	1285	Pd	1004	1004-1506
Na	33.6-50.4	33.6-50.4-134.4	Ag	577-x	577-866
Al	376-564	564-1128	Cd	246-x	246-554
K	18.8	x-37.6	In	236	236-354
V	913-x	1370	Sn	302	226-453
Cr	x-913	x-457	Ta	1072	1206-x
Mn	293-1172	586-1172	W	1733	1733
Fe	913	913-1370	Ir	2007	1338-2007
Ni	913-1370	913-1370	Pt	1338	1338
Cu	845-1266	1266	Au	867	650-867
Zn	305	305-610	Tl	x-253	169-x
As	274-548	274-548-822	Pb	221-331	165-441
Nb	897-1196	1794	Bi	165-331	x-662
Mo	1442	1442-2121	Th	313-626	626-1565
Rh	1442	1442	U	578-1156	419-838

The amount of difference between the two columns of the table should not be surprising. The atomic rotations that determine the azy factors are the same in both cases, but the possible values of these factors have a substantial range of variation, and the influences that affect the values of these factors are not identical. In view of the

participation of the electrons in the resistivity relations, and the large impurity effects, neither of which enters into the volume relations, some difference in the pressures at which the transitions take place can be considered normal. There is, at present, no explanation for those cases in which the internal pressures indicated by the results of the compression and resistance measurements are widely divergent, but differences in the specimens can certainly be suspected.

TABLE 26: *RELATIVE RESISTANCE UNDER COMPRESSION*

Pressure	Calc.	Obs.	Calc.	Obs.	Calc.	Obs.	Calc.	Obs.
(M kg/cm ²)	W 4-8-3		Pt 4-8-2		Rh 4-8-2		Cu 4-8-1½	
0	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
10	.989	.987	.985	.981	.986	.984	.984	.982
20	.977	.975	.971	.963	.973	.968	.969	.965
30	.966	.963	.957	.947	.960	.953	.954	.949
40	.955	.951	.943	.931	.947	.939	.940	.934
50	.945	.940	.929	.916	.934	.925	.925	.920
60	.934	.930	.916	.903	.922	.912	.912	.907
70	.924	.920	.903	.891	.910	.900	.898	.895
80	.914	.911	.890	.880	.897	.8889	.885	.884
90	.904	.903	.878	.870	.886	.880	.872	.875
100	.894	.895	.866	.861	.875	.872	.859	.866
(M kg/cm ²)	Ni 4-8-3 4-8-1½		Fe 4-8-2 4-8-1½		Pd 4-8-2 4-6-3		Zn 4-8-1½ 4-4-2	
0	1000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
10	<u>.978</u>	.982	.978	.977	.980	.980	.938	.937
20	.960	.965	.958	.956	.961	.960	<u>.881</u>	.887
30	.946	.948	.937	.936	.943	.942	.836	.847
40	.933	.933	<u>.918</u>	.919	.925	.925	.810	.812
50	.919	.918	.901	.903	<u>.907</u>	.909	.786	.783
60	.907	.904	.889	.888	.891	.894	.762	.756
70	.894	.892	.875	.875	.880	.881	.740	.733
80	.882	.880	.864	.862	.868	.962	.719	.713
90	.870	.869	.853	.851	.858	.858	.699	.695
100	.858R	.858	.841R	.841	.847R	.847	.679R	.679

Table 26 compares the relative resistances calculated from equation 10-2 with Bridgman's results on some typical elements. The data are presented in the same form as the compressibility tables in Chapter 4, to facilitate comparisons between the two sets of results. This includes showing the azy factors for each element rather than the internal pressures, but the corresponding pressures are available in Table 25. As in the compressibility tables, values above the transition pressures are calculated relative to an observed value as a reference level. The reference value utilized is indicated by the symbol R following the figure given in the "calculated" column.

In those cases where the correct assignment of azy factors and internal pressures above the transition point is not definitely indicated by the corresponding compressibility values, the selections from among the possible values are necessarily based on the empirical measurements, and they are therefore subject to some degree of uncertainty. Agreement between the experimental and the semi-theoretical values in this resistance range therefore validates only the exponential relation in equation 10-2, and does not necessarily confirm the specific values that have been calculated. The theoretical results below the transition points, on the other hand, are quite firm, particularly where the indicated internal pressures are supported by the results of the compressibility measurements. On this basis, the extent of agreement between theory and observation in the values applicable to those elements that maintain the same internal pressures through the full 100.000 kg/cm² pressure range of Bridgman's measurements is an indication of the experimental accuracy. The accuracy thus indicated is consistent with the estimates made earlier on the basis of other criteria.

Inasmuch as the difference in the form of the compressibility equation, $PV^2=k$ (equation 4-4), and that of the pressure-resistance equation, $P^2R=k$ (equation 10-1), is a requirement of the *general* reciprocal relation between space and time specified in the postulates of the Reciprocal System of theory, the joint verification of these two equations is a significant addition to the mass of evidence confirming the validity of this reciprocal relation, the cornerstone of the quantitative expression of the theory of the universe of motion.

As noted earlier, it has been necessary to confine this present discussion of electrical resistance to the resistance of conductors, but a comment on the total range of the resistance values will be appropriate. The resistivity is inversely related to the conductivity in the manner of such reciprocal quantities as speed and energy. Such inverse quantities are symmetrical around the unit value. Resistivity extends from zero to unity. Thus unit resistivity is equal to unit conductivity, and the conductivity then decreases from unity to zero as the equivalent resistivity is increased.

CHAPTER 11

Thermoelectric Properties

As brought out in Chapter 9, the equivalent space in which the thermal motion of the atoms of matter takes place contains a concentration of electrons, the magnitude of which is determined, in the first instance, by factors that are independent of the thermal motion. In the thermal process the atoms move through the electron space as well as through the equivalent of extension space. Where the net time displacement of the atoms of matter provides a time continuum in which the electrons (units of space) can move, a portion of the atomic motion is communicated to the electrons. The thermal motion in the time region environment therefore eventually arrives at an equilibrium between motion of matter through space and motion of space (electrons) through matter.

It should be noted particularly that the motion of the electrons through the matter is a part of the thermal motion, not something separate. A mass m attains a certain temperature T when the effective energy of the thermal motion reaches the corresponding level. It is immaterial from this standpoint whether the energy is that of motion of mass through equivalent space, or motion of space (electrons) through matter, or a combination of the two. In previous discussions of the hypothesis that metallic conduction of heat is due to the movement of electrons, the objection has been raised that there is no indication of any increment in the specific heat due to the thermal energy of the electrons. The development of the Reciprocal System of theory has not only provided a firm theoretical basis for what was previously no more than a hypothesis—the electronic nature of the conduction process—but has also supplied the answer to this objection. The electron movement has no effect on the specific heat because it is not an addition to the thermal motion of the atoms; it is an integral part of the combination motion that determines the magnitude of the specific heat.

Because the factors determining the electron capture from and loss to the environment are independent of the nature of the matter and the amount of thermal motion, the equilibrium concentration is the same in any isolated conductor, irrespective of the material of which the conductor is composed, the temperature, or the pressure. All of these factors do, however, enter into the determination of the thermal energy per electron. Like the gas pressure in a closed container, which depends on the number of molecules and the average energy per molecule, the electric voltage within an isolated conductor is determined by the number of electrons and the average energy per electron. In such a isolated conductor the electron concentration is uniform. The electric voltage is therefore proportional to the thermal energy per electron.

If two conductors of dissimilar composition, copper and zinc, let us say, are brought into contact, the difference in the electron energy level will manifest itself as a voltage differential. A flow of electrons will take place from the conductor with the higher (more negative) voltage, the zinc, to the copper until enough electrons have been transferred to bring the two conductors to the same voltage. What then exists is an equilibrium between a smaller number of relatively high energy electrons in the zinc and a greater number of relatively low energy electrons in the copper.

In this example it is assumed that the voltages in the conductors are allowed to reach an equilibrium. Some more interesting and significant effects are produced where equilibrium is not established. For instance, a continuing current may be passed through the two conductors. If the electron flow is from the zinc to the copper, the electrons leave the zinc with the relatively high voltage that prevails in that conductor. In this case the lower voltage of the electrons in the copper conductor cannot be counterbalanced by an increase in the electron concentration, as all of the electrons that enter the copper under steady flow conditions pass on through. The incoming electrons therefore lose a portion of their energy content in the process of conforming to the new environment. The difference is given up as heat, and the temperature in the vicinity of the zinc-copper junction increases. If the section of the conductor under consideration is part of a circuit in which the electrons return to the zinc, this process is reversed at the copper-zinc junction. Here the energy level of the incoming electrons rises to conform with the

higher voltage of the zinc, and heat is absorbed from the environment to provide the electron energy increment. This phenomenon is known as the *Peltier effect*.

In this Peltier effect a flow of current causes a difference between the temperatures at the two junctions. The *Seebeck effect* is the inverse process. Here a difference in temperature between the two junctions causes a current to flow through the circuit. At the heated junction the increase in thermal energy raises the voltage of the high energy conductor, the zinc, more than that of the low energy conductor, the copper, because the size of the increment is proportional to the total energy. A current therefore flows from the zinc into the copper, and on to the low temperature junction. The result at this junction is the same as in the Peltier effect. The net result is therefore a transfer of heat from the hot junction to the cold junction.

Throughout the discussion in this volume, the term “electric current” refers to the movement of uncharged electrons through conductors, and the term “higher voltage” refers to a greater force, $\frac{1}{2}e^2$, due to a greater concentration of electrons or its equivalent in a greater energy per electron. This electron flow is opposite to the conventional, arbitrarily assigned, “direction of current flow” utilized in most of the literature on current electricity. Ordinarily the findings of this work have been expressed in the customary terms of reference, even though in some cases those findings suggest that an improvement in terminology would be in order. In the present instance, however, it does not appear that any useful purpose would be served by incorporating an unfortunate mistake into an explanation whose primary purpose is to clarify relationships that have been confused by mistakes of other kinds.

A third thermoelectric phenomenon is the *Thomson effect*, which is produced when a current is passed through a conductor in which a temperature gradient exists. The result is a transfer of heat either with or against the temperature gradient. Here the electron energy in the warm section of the conductor is either greater or less than that in the cool section, depending on the thermoelectric characteristics of the conductor material. Let us consider the case in which the energy is greater in the warm section. The electrons that are in thermal equilibrium with the thermally moving matter in this section have a relatively high energy content. These energetic electrons are carried

by the current flow to the cool section of the conductor. Here they must lose energy in order to arrive at a thermal equilibrium with the relatively cold matter of the conductor, and they give up heat to the environment. If the current is reversed, the low energy electrons from the cool section travel to the warm section, where they absorb energy from the environment to attain thermal equilibrium. Both of these processes operate in reverse if the material of the conductor is one of the class of substances in which the effective voltage decreases with an increase in the temperature. There are also some substances in which the response of the voltage to a temperature increment changes direction at some specific temperature level. A similar reversal of the Thomson effect occurs whenever a change of this kind takes place.

The quantitative measure of capability to produce the thermoelectric effects is the *thermoelectric power* of the various conductor materials. This is the electric voltage, expressed either relative to some reference substance, usually lead, or as an absolute value measured against a superconducting material. Neither the theoretical study nor the experimental measurements are far enough advanced to make a quantitative comparison of theory with experimental results feasible at this time, but some of the general considerations that are involved in the quantitative determination can be deduced from theoretical premises.

The basic difference between the thermal motion of the electrons and that of the atoms of matter is in the location of the initial level, or zero point. The zero for the thermal motion of the *atoms* is the equilibrium condition, in which the atom is stationary in a three-dimensional coordinate system of reference because the motion imparted to it by the progression of the natural reference system is counterbalanced by the oppositely directed gravitational motion. On the other hand, the zero for the thermal motion of the *electrons*, the magnitude of the motion of the electrons in the absence of thermal motion, is the natural zero, which, in the context of the stationary reference system, is unit speed, the speed of light. The measure of the energy of the electron motion in matter is the deviation of the speed upward or downward from this unit level.

The fact that the zero energy levels of the positive and negative electron motion are coincident explains why each thermoelectric

effect is a single phenomenon in which the zero level is merely a point in a continuous succession of magnitudes, rather than a discontinuous phenomenon such as the resistance to current flow. The difference between a small positive electron speed and a small negative electron speed is itself relatively small, and within the limits of what can be accomplished by a change in the conditions to which the conductor is subject. Such a change in conditions may therefore reverse the motion. But a substance that is a conductor in one temperature or pressure range does not become an insulator in another range, because the positive zero is the equivalent of the negative infinity, rather than the negative zero, and consequently there is an immense gap between a small positive thermal speed and a small negative speed, in application to the atomic motion. There is, as a consequence, an immense gap between a small positive thermal speed and a small negative speed.

The status of the electron motion as positive or negative is determined by the position that the interacting atom occupies in its rotational group, in the same manner as the effective electric displacement of the atom. Each of these rotational groups consists of two divisions that are positive from the atomic standpoint, followed by two negative divisions. But since the electron is a single rotating system, instead of a double system of the atomic type, the various subdivisions of the atomic series are reduced to half size in application to the electrons. The reversals from positive to negative therefore occur at every divisional boundary in electronic processes, rather than at every second division.

Identification of individual elements as positive or negative from the thermoelectric standpoint is necessarily subject to some qualifications because, as previously mentioned, some elements are positive in one temperature range and negative in another, but a reasonably good test of the theoretical conclusions can be accomplished by comparing the sign of the thermoelectric power as observed at zero degrees C with the divisional status of the elements for which thermoelectric data are available in one of the recent compilations. Table 27 presents such a comparison, omitting the Division I elements of displacements 1 and 2.

TABLE 27: *THERMOELECTRIC POWER*

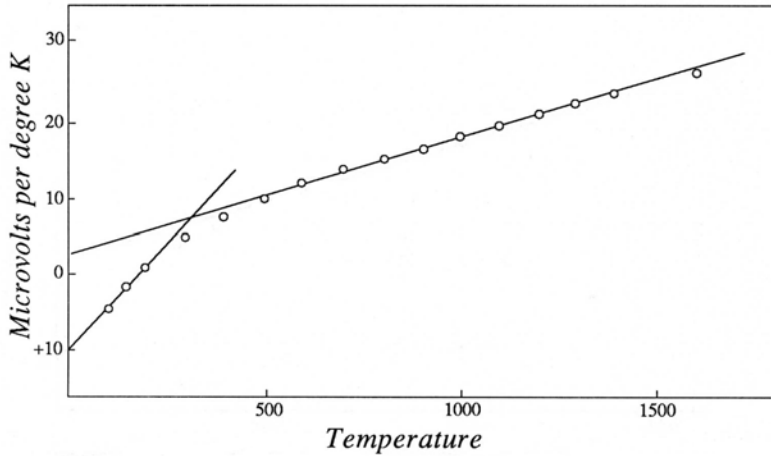
Division				
I	II	III	IV	
Al+	Co-	Cu+	W+	Si-
Ce+	Fe-	Zn+	Ir+	Pb-
	Ni-	Ge+	Pt-	Bi-
	Mo+	Ag+	Au+	
	Pd-	Cd+	Hg-	
		In+	Tl+	
		Sn+		

The reason for the omissions from the tabulation is that here again, as in the resistivity relations discussed in Chapter 10, the first two Division I elements of each rotational group follow a distinctive pattern of their own. In these elements the factor controlling the thermoelectric power is the magnetic rotational displacement, rather than the electric displacement. Because of the single rotation of the electron, the range of magnetic displacements from 1-1 to 4-4 becomes two divisions, with a reversal of sign at the boundaries. For reasons of symmetry, the interior section from 2-2 to 3-3 constitutes one division, in which the displacement one elements, sodium, potassium, and rubidium, have negative thermoelectric voltages. The corresponding members of the outer groups, lithium and cesium, have positive voltages. The displacement two elements may follow either the magnetic or the electric pattern. One of those included in the reference tabulation, calcium, has the same negative voltage as its neighbor, potassium, but magnesium, the corresponding member of the next lower group, takes the positive voltage of the higher Division I elements.

While the theoretical development that is being described in this work has not yet been extended to the quantitative aspects of the thermoelectric effects thus far discussed, it is of interest to note that the relation of the thermoelectric power to temperature has many of the characteristics that we encountered in our previous examination of the response of other properties of matter to temperature changes. This is well illustrated in Fig. 16, which shows the relation between temperature and the absolute thermoelectric power of platinum. Without the captions it would be difficult to distinguish this diagram from one

applicable to thermal expansion, or to the specific heat of an element of one of the lower groups. This is no accident. The curves look alike because the same basic factors are applicable in all of these cases.

FIGURE 16: *ABSOLUTE THERMOELECTRIC POWER-PLATINUM*



In the platinum curve the initial level is positive and the increments due to higher temperature are negative. This behavior is reversed in such elements as tungsten, which has a negative initial level and positive temperature increments up to a temperature of about 1400 K. Above this temperature there is a downward trend. This downward portion of the curve (linear, as usual) is the second segment. At the present stage of the theoretical development it appears probable that a general rule is involved here; that is, the second segment of each curve, the multi-unit segment, is directed toward more negative values, irrespective of the direction of the first (single-unit) segment.

Another thermoelectric effect is the *conduction of heat*. This is a process that is more important from a practical standpoint than those effects that were considered earlier, and it has therefore been given more attention in the present early stage of the development of the theory of the universe of motion. Although the examination of the subject was a somewhat incidental feature of the review of electric current phenomena undertaken in preparation for the new edition of this work, it has produced a fairly complete picture of the heat conductivity of the principal class of conducting metals, together with

a general idea of the manner in which other elements deviate from the general pattern. It was possible to achieve these results in the limited time available because, as it turned out, the metallic conduction of heat is not a complex process, involving difficult concepts such as phonons, orbitals, relaxation processes, electron scattering, and so on, as seen by conventional physics, but a very simple process, capable of being defined by equally simple mathematics, closely related to the mathematical relations governing purely mechanical processes.

In the first situation discussed in this chapter, that in which two previously isolated conductors of different composition are brought into contact, the electron energies in the two conductors are necessarily unequal. As brought out there, the contact results in the establishment of an equilibrium between a larger number of less energetic electrons in one conductor and a smaller number of more energetic electrons in the other. Such an equilibrium cannot be established between two sections of a homogeneous conductor because in this case there is no influence that *requires* either the individual electron energy or the electron concentration to take different values in different locations. If the environmental conditions are uniform, both the energy distribution and the electron concentration attain uniformity throughout the conductor.

However, if one end of a conductor composed of a material such as iron is heated, the energy content of the electrons at that location is increased, and a force differential is generated. Under the influence of the force gradient some of the hot electrons move toward the cold end of the conductor. At that end the newly arrived electrons give up heat in the process of reaching a thermal equilibrium with the atomic motion, and join the concentration of cold electrons previously existing at this location. The resulting higher electron pressure causes a flow of cold electrons back toward the hot end of the conductor. None of the characteristic electrical effects are produced in this process, because the two oppositely directed electron flows are equal in magnitude, and the effects produced by one current are canceled by those produced by the other. The only observable result is a transfer of heat from the hot end of the conductor to the cold end.

It should be noted that no electrostatic potential difference is involved in either of these current flows. This is one of the obstacles

in the way of a simple explanation of heat conduction in the context of conventional physical theory, where electric currents are assumed to result from differences in potential. As explained in Chapter 9, our finding is that all of the forces causing flow of current in the conductor under consideration, that due to the excess energy of the hot electrons, that due to the increased concentration of electrons at the cold end, and that due to electric voltage in general, are forces of a *mechanical* type, not electrostatic forces.

If the material of the conductor is a substance such as copper in which the voltage decreases (becomes less negative) as the temperature rises, the same result is produced in an inverse manner. Here the effective energy of the electrons at the hot end of the conductor is lower than that of the cold electrons. A flow of cold electrons into the hot region therefore takes place. These electrons absorb heat from the environment to attain thermal equilibrium with the matter of the conductor. The resulting increased concentration of hot electrons is then relieved by a flow of some of these electrons back toward the cold end of the conductor. Here, again, the two oppositely directed electron flows produce no net electrical effects.

The *conduction* of heat in metals by movement of *electrons* is essentially the same process as the *convection* of heat by movement of gas or liquid *molecules*. In a closed system, energetic molecules from a hot region move toward a cold region, while a parallel flow carries an equal number of cold molecules back to the hot region. There is only one significant difference between the two heat transfer processes. Because the fluid molecules are subject to a gravitational effect, heat transfer by convection is relatively rapid if it is assisted by a thermally caused difference in density, whereas it is much slower if the diffusion of the hot molecules operates against the gravitational force.

The quantitative measure of the ability of the electron movement to conduct heat is known as the *thermal conductivity*. Its magnitude is determined primarily (perhaps entirely) by the effective specific heat and the temperature coefficient of resistivity, both of which are inversely related to the conductivity. There is a possibility that it may also be affected to a minor degree by some other influences not yet identified, but in any event, all of the modifying influences other than the specific heat are independent of the temperature,

within the range of accuracy of the measurements of the thermal conductivity, and they can be combined into one constant value for each substance. The thermal conductivity of the substance is then this constant divided by the effective specific heat:

$$\text{Thermal conductivity} = \frac{k}{C_p} \quad (11-1)$$

As we saw in the earlier chapters, the specific heat of the conductor materials follows a straight line relation to the temperature in the upper portion of the temperature range of the solid state, and the resistance is linearly related to the temperature at all points. At these higher temperatures, therefore, there is a constant relation between the thermal conductivity and the electrical conductivity (the reciprocal of the resistivity). This relation is known as the *Wiedemann-Franz law*.

The relation expressed in this law breaks down at the lower temperatures, as soon as the specific heat drops below the original straight line. However, the failure of the relation does not occur as soon as would be expected from the normal specific heats of the metals, most of which begin to drop away from the upper linear segment of the curve in the neighborhood of room temperature. The reason for the extension of the high temperature linear relation to a lower temperature in application to thermal conductivity is that the specific heat under the conditions applicable to thermal conduction is not subject to all of the limitations that apply to the transmission of thermal energy by contact between atoms of matter. Instead of going through some intermediate steps, as in the measured specific heats, the effective specific heat in thermal conduction continues on the high temperature basis down to the point where multi-unit motion is no longer possible, and a transition to a single unit basis is mandatory.

The temperature designated as T_0 in the previous discussion, the point at which the specific heat curve reaches the zero level, is the same in thermal conduction as in the atomic contacts, but in the interaction between the electrons and the atoms the single rotating system of the electron adds one half unit to the one unit initial level of the double system of the atom. The initial level of the modified specific heat curve

is therefore $1\frac{1}{2}$ units (-1.98) instead of the usual one unit (-1.32). This makes the slope of the curve somewhat steeper than that of the initial segment of the normal specific heat curve defined in Chapter 5.

The deviation of the thermal conductivity from the constant relation expressed by the Wiedemann-Franz law is the problem with which any theory of thermal conductivity has to deal, and since the explanation derived from the Reciprocal System of theory attributes this deviation to the specific heat pattern, the best way to demonstrate the validity of the explanation appears to be to work backward from the measured thermal conductivities (reference 21), calculate the corresponding theoretical specific heats from equation 11-1, and then compare these calculated specific heats with the theoretical pattern just described.

FIGURE 17: *EFFECTIVE SPECIFIC HEAT IN THERMAL CONDUCTIVITY*

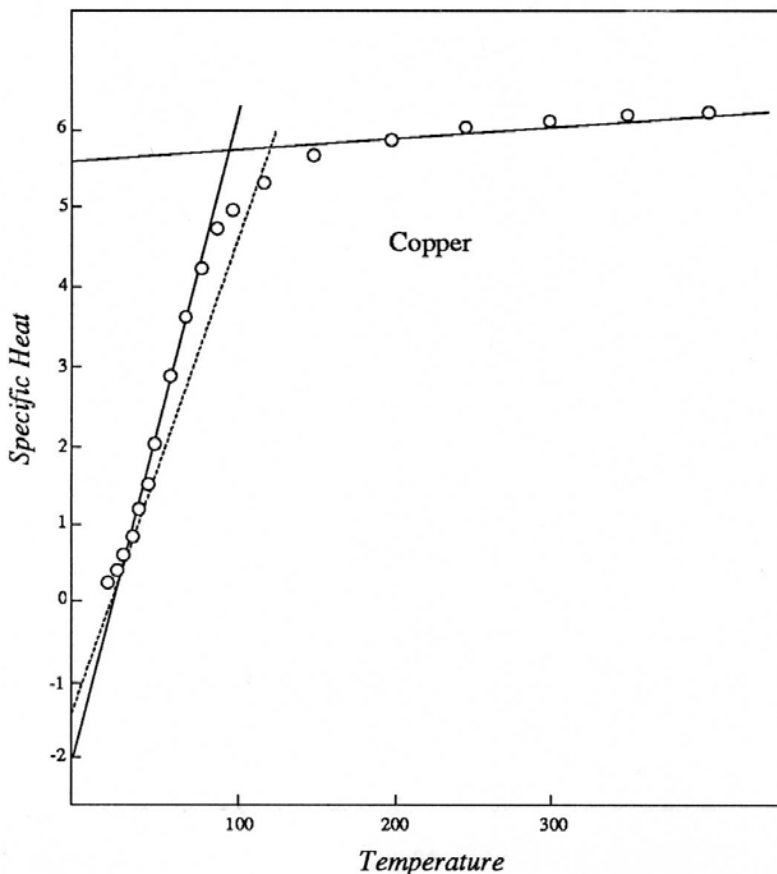
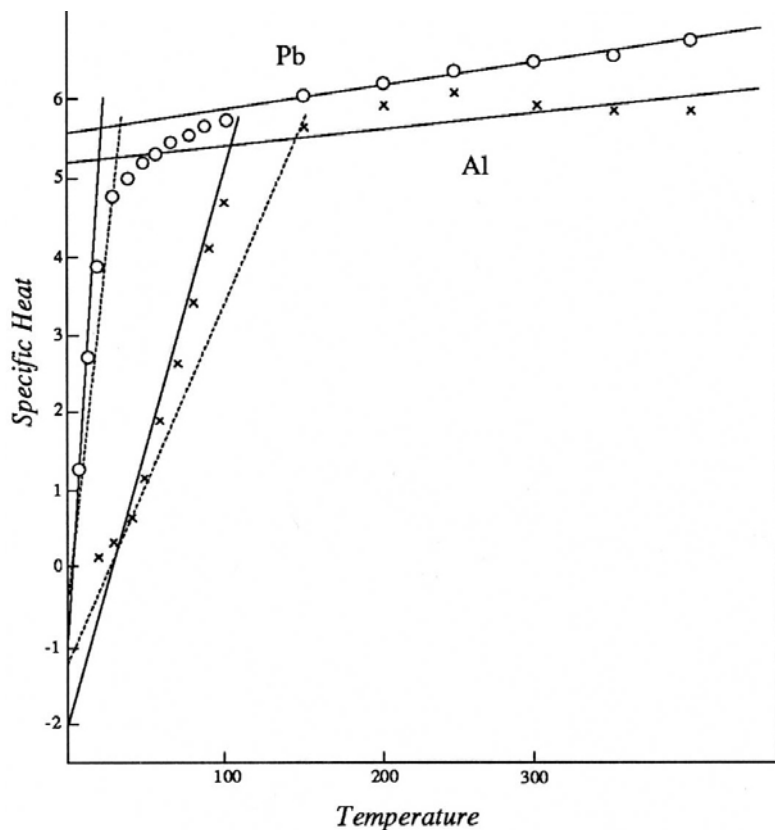


Fig. 17 is a comparison of this kind for the element copper, for which the numerical coefficient of equation 11-1 is 24.0, where thermal conductivities are expressed in watts $\text{cm}^{-2} \text{deg}^{-1}$. The solid lines in this diagram represent the specific heat curve applicable to the thermal conductivity of copper, as defined in the preceding discussion. For comparison, the first segment of the normal specific heat curve of this element is shown as a dashed line. As in the illustrations of specific heat curves in the preceding chapters, the high temperature extension of the upper segment of the curve is omitted in order to make it possible to show the significant features of the curve more clearly. As the diagram indicates, the specific heats calculated from the measured thermal conductivities follow the theoretical lines within the range of the probable experimental errors, except at the lower and upper ends of the first segment, where

FIGURE 18: *EFFECTIVE SPECIFIC HEAT IN THERMAL CONDUCTIVITY*

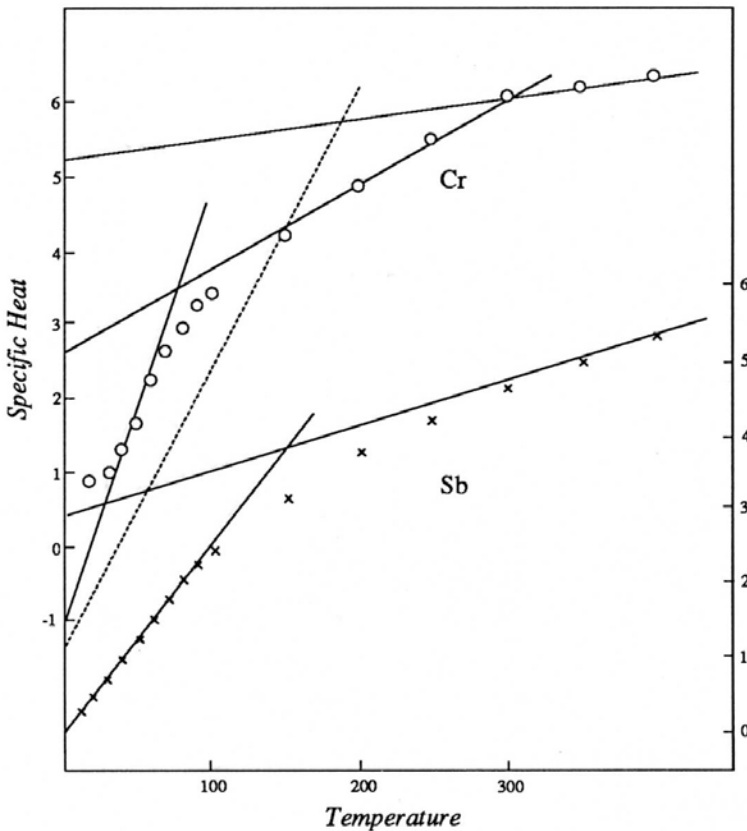


transition curves of the usual kind reflect the deviation of the specific heat of the aggregate from that of the individual atoms.

Similar data for lead and aluminum are presented in Fig. 18.

The pattern followed by the three elements thus far considered may be regarded as the regular behavior, the one to which the largest number of the elements conform. No full scale investigation of the deviations from this basic pattern has yet been undertaken, but an idea of the nature of these deviations can be gained from an examination of the effective specific heat of chromium, Fig. 19. Here the specific heat and temperature values in the low temperature range have only half the usual magnitude. The negative initial specific heat level is -1.00 rather than -2.00 , the temperature of zero specific heat is 16°K rather than 32°K , and the initial level of the upper segment of the curve is 2.62 instead of 5.23 . But

FIGURE 19: *EFFECTIVE SPECIFIC HEAT IN THERMAL CONDUCTIVITY*



this upper segment of the modified curve intersects the upper segment of the normal curve at the Neel point, 311 K, and above this temperature the effective specific heat of chromium in thermal conductivity follows the regular specific heat pattern as defined in Chapter 5.

Another kind of deviation from the regular pattern is seen in the curve for antimony, also shown in Fig. 19. Here the initial level of the first segment is zero instead of the usual negative value. The initial level of the second segment is the half-size value 2.62. Antimony thus combines the two types of deviation that have been mentioned.

As indicated earlier, it has not yet been determined whether any factors other than the resistivity coefficient enter into the constant k of equation 11-1. Resolution of this issue is complicated by the wide margin of uncertainty in the thermal conductivity measurements. The authors of the compilation from which the data used in this work were taken estimate that these values are correct only to within 5 to 10 percent in the greater part of the temperature range, with some uncertainties as high as 15 percent. However, the agreement between the plotted points in Figures 17, 18, and 19, and the corresponding theoretical curves shows that most of the data represented in these diagrams are more accurate than the foregoing estimates would indicate, except for the aluminum values in the range from 200 to 300 °K.

In any event, we find that for the majority of the elements included in our preliminary examination, the product of the empirical value of the factor k in equation 11-1 and the temperature coefficient of resistivity is between 0.14 and 0.18. Included are the best known and most thoroughly studied elements, copper, iron, aluminum, silver, etc., and a range of k values extending all the way from the 25.8 of silver to 1.1 in antimony. This rather strongly suggests that when all of the disturbing influences such as impurity effects are removed, the empirical factor k in equation 11.1 can be replaced by a purely theoretical value k_r , in which a theoretically derived conversion constant, k , in the neighborhood of $0.15 \text{ watts cm}^{-2} \text{ deg}^{-1}$ is divided by a theoretically derived coefficient of resistivity.

The impurity effects that account for much of the uncertainty in the general run of thermal conductivity measurements are still more

prominent at very low temperatures. At least on first consideration, the theoretical development appears to indicate that the thermal conductivity should follow the same kind of a probability curve in the region just above zero temperature as the properties discussed in the preceding chapters. In many cases, however, the measurements show a minimum in the conductivity at some very low temperature, with a rising trend below this level. On the other hand, some of the elements that are available in an extremely pure state show little or no effect of this kind, and follow curves similar to those encountered in the same temperature range during the study of other properties. It is not unlikely that this will prove to be the general rule when more specimens are available in a pure enough state. It should be noted that an ordinary high degree of purity is not enough. As the data compilers point out, the thermal conductivities in this very low temperature region are "highly sensitive to small physical and chemical variations of the specimens."

CHAPTER 12

Scalar Motion

It was recognized from the beginning of the development of the theory of the universe of motion that the basic motions are necessarily scalar. This was stated specifically in the first published description of the theory, the original (1959) edition of *The Structure of the Physical Universe*. It was further recognized, and emphasized in that 1959 publication, that the rotational motion of the atoms of matter is one of these basic scalar motions, and therefore has an inward translational effect, which we can identify as gravitation. Throughout the early stages of the theoretical development, however, there was some question as to the exact status of rotation in a system of scalar motions, inasmuch as rotation, as ordinarily conceived, is directional, whereas scalar quantities, by definition, have no directions. At first this issue was not critical, but as the development of the theory was extended into additional physical areas, more types of motion of a rotational character were encountered, and it became necessary to clarify the nature of scalar rotation. A full scale investigation of the subject was therefore undertaken, the results of which were reported in *The Neglected Facts of Science*, published in 1982.

The existence of scalar motion is not recognized by present-day physics. Indeed, motion is usually defined in such a way that scalar motion is specifically excluded. This type of motion enters into observable physical phenomena in a rather unobtrusive manner, and it is not particularly surprising that its existence remained unrecognized for a long time. However, a quarter of a century has elapsed since that existence was brought to the attention of the scientific community in the first published description of the universe of motion, and it is hard to understand why so many individuals still seem unable to recognize that there are several observable types of motion that cannot be other than scalar.

For instance, the astronomers tell us that the distant galaxies are all moving radially outward away from each other. The full significance of this galactic motion is not apparent on casual consideration, as we see each of the distant galaxies moving outward from our own location, and we are able to locate each of the observed motions in our spatial reference system in the same manner as the familiar motions of our everyday experience. But the true character of this motion becomes apparent when we examine the relation of our Milky Way galaxy to this system of motions. Unless we take the stand that our galaxy is the only stationary object in the universe, an assumption that few scientists care to defend in this modern era, we must recognize that our galaxy is moving away from all of the others; that is, it is moving in all directions. And since it is conceded that our galaxy is not unique, it follows that *all* of the widely separated galaxies are moving outward in all directions. Such a motion, which takes place uniformly in all directions has no *specific* direction. It is completely defined by a magnitude (positive or negative), and is therefore *scalar*.

A close examination of gravitation shows that the gravitational motion is likewise scalar, differing from the motion of the galaxies only in that it is negative (inward) rather than positive (outward). The resemblance to the motion of the galaxies can easily be seen if we consider a system of gravitating objects isolated in space—perhaps a group of galaxies relatively close to each other. From our knowledge of the gravitational effects we can deduce that each of these objects will move inward toward all of the others. Here again the motion is scalar. Each object is moving inward in all directions.

A small-scale example of the same kind of motion can be seen in the motion of spots on the surface of an expanding balloon, often used as an analogy by those who undertake to explain the nature of the motions of the distant galaxies. Here, too, each individual is moving outward from all others. If the expansion is terminated, and succeeded by a contraction, the motions are reversed, and each spot then moves inward toward all others, as in the gravitational motion.

In the case of the expanding balloon there is a known physical mechanism that is causing the expansion, and our understanding of

this mechanism makes it evident that all *locations* on the balloon surface are moving. The spots on this surface have no motion of their own. They are merely being carried along by the movement of the locations that they occupy. According to the astronomers' current view, the recession of the distant galaxies is the same kind of a process. As Paul Davies explains:

Many people (including some scientists) think of the recession of the galaxies as due to the explosion of a lump of matter into a pre-existing void, with the galaxies as fragments rushing through space. This is quite wrong... The expanding universe is not the motion of the galaxies through space away from some centre, but is the steady expansion of space.²²

Here, again, it is the *locations* that are moving, carrying the galaxies along with them. But in this case there is no known physical mechanism to account for the movement. Like the expansion of the balloon, the “steady expansion of space” is merely a description, not an explanation, of the movement. All that the observations tell us is that an outward scalar motion of physical locations is taking place, carrying the galaxies with it.

The postulates of the Reciprocal System of theory, the theory of the universe of motion, generalize this type of motion. They define a universe in which scalar motion of physical locations is the basic form of motion from which all physical entities and phenomena are derived. The manner in which this type of motion manifests itself to observation therefore has an important bearing on the nature of fundamental physical phenomena.

This situation is a good example of the way in which important information is often overlooked because no one spends the time and effort that are required in order to make a thorough study of a seemingly unimportant observation. It has long been recognized that the motion of spots on the surface of an expanding balloon is, in some way, different from the ordinary motions of our everyday experience. The mere fact that this balloon motion is so widely used as an analogy in explaining the recession of the distant galaxies is clear evidence of this general recognition. But the galaxies seem to be a special case, and expanding balloons do not play any significant

part in normal physical activity. Consequently, no one has been much interested in the physics of these objects, and this admittedly unique kind of motion was never subjected to a critical examination prior to the investigation of scalar motion that was undertaken in the course of the theoretical development reported in the several volumes of this work. The finding that the fundamental motion of the universe is scalar revolutionizes this situation. The motions of the galaxies, gravitating objects, and spots on the surface of an expanding balloon are obviously the kind of motions—scalar motions—that our theory identifies as fundamental.

Scientists are usually, with ample justification, reluctant to accept a hypothesis that postulates the existence of phenomena that are unknown to observation. It should therefore be emphasized that scalar motion is not an *unobserved* phenomenon; it is an observed phenomenon that has not heretofore been *recognized* in its true character. Once the motions identified in the foregoing paragraphs have been critically examined, and their scalar character has been recognized, the *existence of scalar motion* is no longer a hypothesis; it is a demonstrated physical fact. The existence of *other* scalar motions, as required by the theory of the universe of motion is then a natural and logical corollary, and those observed phenomena that have the theoretical properties of scalar motion can legitimately be identified as scalar motions.

A one-dimensional scalar motion of a physical location is defined by a magnitude, and can therefore be represented one-dimensionally as a point, or an assemblage of points, moving along a straight line. Introduction of a *reference point*—that is, coupling the motion to the reference system at a specific point in that system—enables distinguishing between positive motion, outward from the reference point, and negative motion, inward toward the reference point. The direction imputed to the motion may be a *constant* direction, as in the case of the translational motion of the photon, the direction of which is determined by chance at the time of emission, unless external factors intervene. The key point disclosed by our investigation is that the direction is not *necessarily* constant. A discontinuous, or non-uniform, change of direction could be maintained only by a repeated application of an external force, but it has been known

from the time of Galileo that a continuous and uniform *change* of position or direction is just as permanent and just as self-sustaining as a condition of rest. Our finding merely extends this principle to the assignment of direction to scalar motion.

As an illustration, let us consider the motion of point X, originating at point A, and initially proceeding in the direction AB in three-dimensional space. Then let us assume that line AB is rotated around an axis perpendicular to it, and passing through point A. This does not change the inherent nature or magnitude of the motion of point X, which is still moving radially outward from point A at the same speed as before. What has been changed is the direction of the movement, which is not a property of the motion itself, but a feature of the relation between the motion and three-dimensional space. Instead of continuing to move outward from A in the direction AB, point X now moves outward in *all* directions in the plane of rotation. If that plane is then rotated around another perpendicular axis, the outward motion of point X is distributed over all directions in space. It is then a *rotationally distributed scalar motion*.

The results of such a distributed scalar motion are totally different from those produced by a combination of vectorial motions in different directions. The combined effects of the magnitudes and directions of vectorial motions can be expressed as vectors. The results of addition of these vectors are highly sensitive to the effects of direction. For example, a vectorial motion AB added to a vectorial motion AB' of equal magnitude, but diametrically opposite direction, produces a zero resultant. Similarly, vectorial motions of equal magnitude in all directions from a given point add up to zero. But a scalar motion retains the same positive (outward) or negative (inward) magnitude regardless of the manner in which it is directionally distributed.

None of the types of scalar motion that have been identified can be represented in a fixed spatial reference system in its true character. Such a reference system cannot represent simultaneous motion in all directions. Indeed, it cannot represent motion in more than *one* direction. In order to represent a system of two or more scalar motions in a spatial reference system it is necessary to define a reference point for the system as a whole; that is, the scalar system must be coupled to the reference system in such a way that one of

the moving locations in the scalar system is arbitrarily defined as motionless (from the scalar standpoint) relative to the reference system. The direction imputed to the motion of each of the other objects, or physical locations, in the scalar system is then its direction relative to the reference point.

For example, if we denote our galaxy as A, the direction of the motion of distant galaxy X, as we see it, is AX. But observers in galaxy B, if there are any, see galaxy X as moving in the different direction BX, those in galaxy C see the direction as CX, and so on. The significance of this dependence of the direction on the reference point can be appreciated when it is contrasted with the corresponding aspect of vectorial motion. If an object X is moving vectorially in the direction AX when viewed from location A, it is also moving in this same direction AX when viewed from any other location in the reference system.

It should be understood that the immobilization of the reference point in the reference system applies only to the representation of the *scalar* motion. There is nothing to prevent an object located at the reference point, the reference object we may call it, from acquiring an *additional* motion of a vectorial character. For example, the expanding balloon may be resting on the floor of a moving vehicle, in which case the reference point is in motion vectorially. Where an additional motion of this nature exists, it is subject to the same considerations as any other vectorial motion.

The coupling of a system of scalar motions to a fixed reference system at a reference point does not alter the rate of separation of the members of the scalar system. The arbitrary designation of the reference point as motionless (from the scalar standpoint) therefore makes it necessary to attribute the motion of the reference point, or object, to the other points or objects in the system.

This conclusion that the observed change of position of an object B is due, in part, to the motion of some different object A may be hard for those who are thinking in terms of the conventional view of the nature of motion to accept, but it can easily be verified by consideration of a specific example. Any two spots on the surface of an expanding balloon, for instance, are moving away from each other; that is, they

are *both* moving. While spot X moves away from spot Y, spot Y is coincidentally moving away from spot X. Placing the balloon in a reference system does not alter these motions. The balloon continues expanding in exactly the same way as before. The distance XY continues to increase at the same rate, but if X is the reference point, it is *motionless in the reference system* (so far as the scalar motion is concerned), and the entire increase in the distance XY, including that due to the motion of X, has to be attributed to the motion of Y.

The same is true of the motions of the distant galaxies. The recession that is measured is merely the increase in distance between our galaxy and the one that is receding from us. It follows that a part of the increase in separation that we attribute to the recession of the other galaxy is actually due to motion of our own galaxy. This is not difficult to understand when, as in the case of the galaxies, the reason why objects appear to move faster than they actually do is obviously the arbitrary assumption that our location is stationary. What is now needed is a recognition that this is a general proposition. The same result follows whenever a moving object is arbitrarily taken as stationary for reference purposes. The motion of *any* reference point of a scalar motion is seen, *by the reference system*, in the same way in which we view our motion in the galactic system; that is, the motion that is frozen by the reference system is seen as motion of the distant objects.

This transfer of motion from one object to another by reason of the manner in which scalar motion is represented in the reference system has no significant consequences in the galactic situation, as it makes no particular difference to us whether galaxy X is receding from us, or we are receding from it, or both. But the questions as to which objects are actually moving, and how much they are moving, have an important bearing on other scalar motions, such as gravitation. With the benefit of the information now available, it is evident that the rotation of the atoms of matter described in Volume I is a rotationally distributed negative (inward) scalar motion. By virtue of that motion, each atom, irrespective of how it may be moving, or not moving, vectorially, is moving inward toward all other atoms of matter. This inward motion can obviously be identified as gravitation. Here, then, we have the answer to the question as to the origin of

gravitation. The same thing that makes an atom an atom—the scalar rotation—causes it to gravitate.

Although Newton specifically disclaimed making any inference as to the mechanism of gravitation, the fact that there is no time term in his equation implies that the gravitational effect is instantaneous. This, in turn, leads to the conclusion that gravitation is “action at a distance,” a process in which one mass acts upon another distant mass without an intervening connection. There is no experimental or observational evidence contradicting the instantaneous action. As noted in Volume I, even in astronomy, where it might be presumed that any inaccuracy would be serious, in view of the great magnitudes involved, “Newtonian theory is still employed almost exclusively to calculate the motions of celestial bodies.”²³

However, instantaneous action at a distance is philosophically unacceptable to most physicists, and they are willing to go to almost any lengths to avoid conceding its existence. The hypothesis of transmission through a “luminiferous ether” served this purpose when it was first proposed, but as further studies were made, it became obvious that no physical substance could have the contradictory properties that were required of this hypothetical medium.

Einstein’s solution was to abandon the concept of the ether as a “substance”—something physical—and to introduce the idea of a quasi-physical entity, a phantom medium that is assumed to have the capabilities of a physical medium without those limitations that are imposed by physical existence. He identifies this phantom medium with space, but concedes that the difference between his space and the ether is mainly semantic. He explains, “We shall say our space has the physical property of transmitting waves, and so omit the use of a word [ether] we have decided to avoid.”²⁴ Since this space (or ether) must exert physical effects without being physical, Einstein has difficulty defining its relation to physical reality. At one time he asserts that “according to the general theory of relativity space is endowed with physical qualities,”²⁵ while in another connection he says that “The ether of the general theory of relativity [which he identifies as space] is a medium which is itself devoid of all mechanical and kinematical qualities.”²⁶ Elsewhere, in a more

candid statement, he concedes, in effect, that his explanations are not persuasive, and advises us just to “take for granted the fact that space has the physical property of transmitting electromagnetic waves, and not to bother too much about the meaning of this statement.”²⁷

Einstein’s successors have added another dimension to the confusion of ideas by retaining this concept of space as quasi-physical, something that can be “curved” or otherwise manipulated by physical influences, but transferring the “ether-like” functions of Einstein’s “space” to “fields.” According to this more recent view, matter exerts a gravitational effect that creates a gravitational field, this field transmits the effect at the speed of light, and finally the field acts upon the distant object. Various other fields—electric, magnetic, etc.—are presumed to coexist with the gravitational field, and to act in a similar manner.

The present-day “field” is just as intangible as Einstein’s “space.” There is no physical evidence of its existence. All that we know is that if a test object of an appropriate type is placed within a certain region, it experiences a force whose magnitude can be correlated with the distance to the location of the originating object. What existed *before* the test object was introduced is wholly speculative. Faraday’s hypothesis was that the field is a condition of stress in the ether. Present-day physicists have transferred the stress to space in order to be able to discard the ether, a change that has little identifiable meaning. As R. H. Dicke puts it, “One suspects that, with empty space having so many properties, all that had been accomplished in destroying the ether was a semantic trick. The ether had been renamed the vacuum.”²⁸ P. W. Bridgman, who reviewed this situation in considerable detail, arrived at a similar conclusion. The results of analysis, he says, “suggest that the role played by the field concept is that of an intellectual dummy, which cancels out of the final result.”²⁹

The theory of the universe of motion gives us a totally different view of this situation. In this universe the reality is motion. Space and time have a real existence only where, and to the extent that, they actually exist as components of motion. On this basis extension space, the space that is represented by the conventional reference system, is no more than a frame of reference for the spatial

magnitudes and directions of the entity, motion, that actually exists. It follows that extension space cannot have any physical properties. It cannot be “curved” or modified in any other way by physical means. Of course, the reference system, being nothing but a human contrivance, could be altered conceptually, but such a change would have no physical significance.

The status of extension space as a purely mental concept devised for reference purposes, rather than a physical entity, likewise means that this space is not a container, or background, for the physical activity of the universe, as assumed by conventional science. In that conventional view, everything physically real is contained within the space and time of the spatio-temporal reference system. When it becomes necessary to postulate something outside these limits in order to meet the demands of theory construction, it is assumed that such phenomena are, in some way, unreal. As Werner Heisenberg puts it, they do not “exist objectively in the same sense as stones or trees exist.”³⁰

The development of the theory of the universe of motion now shows that the conventional spatio-temporal system of reference does not contain everything that is physically real. On the contrary, it is an *incomplete* system that is not capable of representing the full range of motions which exist in the physical universe. It cannot represent motion in more than one scalar dimension; it cannot represent a scalar *system* in which all elements are moving; nor can it correctly represent the position of an individual object that is moving in all directions simultaneously (that is, an object whose motion is scalar, and therefore has no specific direction). Many of the other shortcomings of this reference system will not become apparent until we examine the effects of very high speed motion in Volume III, but those that have been mentioned have a significant impact on the phenomena that we are now examining.

The inability to represent more than one dimension of scalar motion is a particularly serious deficiency, inasmuch as the postulated three-dimensionality of the *universe of motion* necessarily permits the existence of three dimensions of *motion*. Only one dimension of *vectorial* motion is possible, because all three dimensions of space are

required in order to represent the directions of this one-dimensional motion, but *scalar* motion has magnitude only, and a three-dimensional universe can accommodate scalar motion in all three of its dimensions.

Since the conventional reference system cannot represent all of the distributed scalar motions, and present-day science does not recognize the existence of any motions that cannot be represented in that system, it has been necessary for the theorists to make some arbitrary assumptions as a means of compensating for the distortion of the physical picture due to this deficiency of the reference system. One of the principal steps taken in this direction is the introduction of the concept of “fundamental forces,” autonomous entities that exist in their own right, and not as properties of something more basic. The present tendency is to regard these so-called fundamental forces as the sources of all physical activity, and the currently popular goal of the theoretical physicists, the formulation of a “grand unified theory,” is limited to finding a common denominator of these forces.

Gravitation is, in a way, an exception, as the currently popular hypothesis as to the nature of the gravitational force, Einstein’s general theory of relativity, does attempt an explanation of its origin. According to this theory, the gravitational force is due to a distortion of space resulting from the presence of matter. So far as can be determined from the scientific literature, no one has the slightest idea as to how such a distortion of space could be accomplished. Arthur Eddington expressed the casual attitude of the scientific community toward this issue in the following statement: “We do not ask how mass gets a grip on space-time and causes the curvature which our theory postulates.”³¹ But unless the question is asked, the answer is not forthcoming. In Newton’s theory the gravitational force originates from mass in a totally unexplained manner. In Einstein’s theory it is a result of a distortion, or “curvature,” of space that is produced by mass in a totally unexplained manner. Thus, whatever its other merits may be, the current theory (general relativity) accomplishes no more toward accounting for the *origin* of the gravitational force than its predecessor.

In order to arrive at such an explanation we need to recognize that force is not an autonomous entity; it is a *property of motion*. The

motion of an individual mass unit is measured in terms of speed (or velocity). The total amount of motion in a material aggregate is then the product of the speed and the number of mass units, a quantity formerly called “quantity of motion,” but now known as momentum. The rate of change of the motion of the individual unit is acceleration; that of the total quantity of motion is force. The force is thus the total quantity of acceleration.

The significance of this, in the present connection, is that force not only *produces* an acceleration when applied to a mass (a fact that is currently recognized), it is an acceleration *prior to that application* (a fact that is currently overlooked or disregarded). In other words, the acceleration is simply transferred. For example, when a rocket is fired, the total “quantity of acceleration” available for application to the rocket (the force) is the sum of the quantities of acceleration of the individual particles of the gas produced from the propellant. The division of this total quantity among the mass units of the rocket determines the acceleration of each individual unit, and therefore of the rocket as a whole.

Since force is a property of a motion rather than an autonomous entity, it follows that wherever there is a force there must also be a motion of which the force is a property. This leads to the conclusion that a gravitational force field is a region of space in which gravitational motions exist. In the context of conventional physical thought this conclusion is unacceptable, since there are no moving entities in an unoccupied field.

The information about the nature of scalar motion developed in the earlier pages clarifies this situation. A material aggregate is moving gravitationally in all directions, but the conventional spatial reference system is unable to represent a system of motions of this nature in its true character. As previously indicated, where the scalar motion AB of object A (the massive object now under consideration) toward object B (the test mass) cannot be represented in the reference system because of the limitations of that system, this motion AB is shown as a motion BA; that is, a motion of the test mass B toward the massive object A, constituting an addition to the actual motion of that test mass. Because of the spherical distribution of the scalar motions of

the atoms of mass A, the magnitude of the motion imputed to mass B depends on its distance from A, and is inversely proportional to the square of that distance. Thus each point in the region surrounding A corresponds to a specific fraction of the motion of A, representing the amount of motion that *would be* imputed to a unit mass, *if that mass is actually placed at this particular point*.

Here, then, is the explanation of the gravitational field (and, by extension, all other fields of the same nature). The field is not something physically real in the space; “for the modern physicist as real as the chair on which he sits,”³² as asserted by Einstein. Nor is it, as Faraday surmised, a stress in the ether. Neither is it some kind of a change in the properties of space, as envisioned by present-day theorists. It is simply the pattern of the magnitudes of the motions of one mass that have to be imputed to other masses because of the inability of the reference system to represent scalar motion as it actually exists.

No doubt this assertion that what appears to be a motion of one object is actually, in large part, a motion of a different object is somewhat confusing to those who are accustomed to conventional ideas about motion. But once it is realized that scalar motion exists, and because it has no inherent direction it may be distributed over all directions, it is evident that the reference system cannot represent this motion in its true character. In the preceding analysis we have determined just how the reference system *does* represent this motion that it cannot represent correctly.

This may appear to be a return to the action at a distance that is so distasteful to most scientists, but, in fact, the apparent action on distant objects is an illusion created by the introduction of the concept of autonomous forces to compensate for the shortcomings of the reference system. If the reference system were capable of representing all of the scalar motions in their true character, there would be no problem. Each mass would then be seen to be pursuing its own course, moving inward in space independently of other objects.

In this case, accepted scientific theory has gone wrong because prejudice supported by abstract theory has been allowed to override the results of physical observations. The observers keep calling attention to the absence of evidence of the finite propagation time

that current theory ascribes to the gravitational effect, as in this extract from a news report of a conference at which the subject was discussed:

When it [the distance] is astronomical, the difficulty arises that the intermediaries need a measurable time to cross, while the forces in fact seem to appear instantaneously.³³

But it is assumed that we must accept either a finite propagation time or action at a distance, which, as Bridgman once said, is “a concept to which many physicists have a violent allergy.”³⁴ Einstein’s theory, which supports the propagation hypothesis, has therefore been accorded a status superior to the observations. The following statement from a physicist brings this point out explicitly:

Nowadays we are also convinced that gravitation progresses with the speed of light. This conviction, however, does not stem from a new experiment or a new observation, it is a result solely of the theory of relativity.³⁵

This is another example of a practice that has been the subject of critical comment in several different connections in the preceding pages of this and the earlier volume. Overconfidence in the existing body of scientific knowledge has led the investigators to assume that all alternatives in a given situation have been considered. It is then concluded that an obviously flawed hypothesis must be accepted, in spite of its shortcomings, because “there is no other way.” Time and again in the earlier pages, development of the theory of the universe of motion has shown that there is another way, one that is free from the objectionable features. So it is in this case. It is not necessary either to contradict observation by assuming a finite speed of propagation or to accept action at a distance.

Some of the most significant consequences of the existence of scalar motion are related to its *dimensions*. This term is used in several different senses, two of which are utilized extensively in this work. When physical quantities are resolved into component quantities of a fundamental nature, these component quantities are called dimensions. Identification of the dimensions, in this sense, of the basic physical quantities has been an important feature of the

development of theory in the preceding pages. In a different sense of the term, it is generally recognized that space is three-dimensional.

Conventional physics recognizes motion in three-dimensional space, and represents motions of this nature by lines in a three-dimensional spatial coordinate system. But these motions which exist in three dimensions of space are only one-dimensional motions. Each individual motion of this kind can be characterized by a vector, and the resultant of any number of these vectors is a one-dimensional motion defined by the vector sum. All three dimensions of the spatial reference system are required for the representation of one-dimensional motion, and there is no way by which the system can indicate a change of position in any other dimension. However, the postulate that the universe of motion is three-dimensional carries with it the existence of three dimensions of *motion*. Thus there are two dimensions of motion in the physical universe *that cannot be represented in the conventional spatial reference system*.

In common usage the word “dimensions” is taken to mean spatial dimensions, and reference to three dimensions is ordinarily interpreted geometrically. It should be realized, however, that the geometric pattern is merely a graphical representation of the relevant physical magnitudes and directions. From the mathematical standpoint an *n*-dimensional quantity is one that requires *n* independent magnitudes for a complete definition. Thus a scalar motion in three dimensions, the maximum in a three-dimensional universe, is defined in terms of three independent magnitudes. One of these magnitudes—that is, the magnitude of one of the dimensions of scalar motion—can be further divided dimensionally by the introduction of *directions* relative to a spatial reference system. This expedient resolves the one-dimensional scalar magnitude into three orthogonally related sub-magnitudes, which together with the directions, constitute vectors. But no more than one of the three scalar magnitudes that define a three-dimensional scalar motion can be subdivided vectorially in this manner.

Here is a place where recognition of the existence of scalar motion changes the physical picture radically. As long as motion is viewed entirely in vectorial terms—that is, as a change of position relative to

the spatial reference system—there can be no motion other than that represented in that system. But since scalar motion has magnitude only, there can be motion of this character in *all three* of the existing dimensions of the physical universe. It should be emphasized that these dimensions of scalar motion are *mathematical* dimensions. They can be represented geometrically only in part, because of the limitations of geometrical representation. In order to distinguish these mathematical dimensions of motion from the geometric dimensions of space in which one dimension of the motion takes place, we are using the term “scalar dimension” in a manner analogous to the use of the term “scalar direction” in the earlier pages of this and the preceding volume.

CHAPTER 13

Electric Charges

The history of the development of a mathematical understanding of electricity and magnetism has been one of the great success stories of science and engineering. With the benefit of this information, a type of phenomena totally unknown up to a few centuries ago has been harnessed in a manner that has revolutionized life in the more advanced human societies. But in a strange contrast, this remarkable record of success in the identification and application of the mathematical relations involved in these phenomena coexists with an almost complete lack of understanding of the basic nature of the quantities with which the mathematical expressions are dealing.

In order to have a reasonably good conceptual understanding of electricity and magnetism, we need to be able to answer questions such as these:

What is an electric charge?

What is magnetism?

What is an electric current?

What is an electric field?

What is mass?

What is the relation between mass and charge?

How are electric and magnetic forces produced?

How do they differ from the gravitational force?

How are these forces transmitted?

What is the reason for the direction of the electromagnetic force?

Why do masses interact only with masses, charges with charges?

How are charges induced in electrically neutral objects?

Conventional science has no answers for most of these questions. To rationalize the failure to discover the explanations, the physicists tell us that we should not ask the questions:

The question “What is electricity?”—so often asked—is... meaningless.³⁶ (E. N. da C. Andrade)

What is electricity?... Definitions that cannot, in the nature of the case, be given, should not be demanded.³⁷ (Rudolf Carnap)

The difficulty in accounting for the origin of the basic forces is likewise evaded. It is observed that matter exerts a gravitational force, an electric charge exerts an electric force, and so on, but the theorists have been unable to identify the origin of these forces. Their reaction has been to evade the issue by characterizing the forces as autonomous, “fundamental conceptions of physics” that have to be taken as given features of the universe. These forces are then assumed to be the original sources of all physical activity.

So far as anyone knows at present, all events that take place in the universe are governed by four fundamental types of forces.³⁸

As pointed out in Chapter 12, this assumption is obviously invalid, as it is in direct conflict with the accepted *definition* of force. But those who are desperately anxious to have *some* kind of a theory of the phenomena that are involved close their eyes to this conflict.

After having “solved” the problem of the origin of the forces by assuming it out of existence, the theorists have proceeded to solve the problem of the transmission of the basic forces in a similar manner. Since they have no explanation for this phenomenon, they provide a substitute for an explanation by equating this transmission with a different kind of phenomenon for which they believe they have at least a partial explanation. Electromagnetic radiation has both electric and magnetic aspects, and is unquestionably a transmission process. In their critical need for some kind of an explanation of the transmission of electric and magnetic forces, the theory constructors have seized on this tenuous connection, and have assumed that electromagnetic radiation is the carrier of the electrostatic and magnetostatic forces. Then, since the gravitational force is clearly analogous to those two forces, and can be represented by the same kind of mathematical expressions, it has been further assumed that some sort of gravitational radiation must also exist.

But there is ample evidence to show that these forces are *not* transmitted by radiation. As brought out in Volume I, gravitation and radiation are

processes of a totally different kind. Radiation is an energy transmission process. A quantity of radiant energy is produced in the form of photons. The movement of these photons then carries the energy from the point of origin to a destination, where it is delivered to the receiving object. No movement of either the originating object or the receiving object is required. At either end of the path the energy is recognizable as such, and is readily interchangeable with other forms of energy.

Gravitation, on the other hand, is *not* an energy transmission process. The (apparent) gravitational action of one mass upon another does not alter the total external energy content (potential plus kinetic) of either mass. Each mass that moves in response to the gravitational force acquires a certain amount of kinetic energy, but its potential energy is decreased by the same amount, leaving the total unchanged. As stated in Volume I, gravitational, or potential, energy is purely an energy of position: that is, for any specific masses, the mutual potential energy is determined entirely by their spatial separation.

All that has been said about gravitation is equally applicable to electrostatics and magnetostatics. Each member of any system of two or more objects (apparently) interacting electrically or magnetically has a potential energy determined by the magnitudes of the charges and the intervening distance. As in the gravitational situation, if the separation between the objects is altered by reason of the static forces, an increment of kinetic energy is imparted to one or more of the objects. But its, or their, potential energy is decreased by the same amount, leaving the total unchanged. This is altogether different from a process such as electromagnetic radiation which carries energy *from* one location *to* another. Energy of position in space cannot be propagated in space. The concept of transmitting this kind of energy from one spatial position to another is totally incompatible with the fact that the magnitude of the energy is *determined by* the spatial separation.

As stated earlier, the coexistence of an almost total lack of conceptual understanding of electric and magnetic fundamentals with a fully developed system of mathematical relations and representations seems incongruous. In fact, however, this is the normal initial result of the manner in which scientific investigation is usually handled.

A complete theory of any physical phenomenon consists of two distinct components, a mathematical formulation and a conceptual structure, which are largely independent. In order to constitute a complete and accurate definition of the phenomenon, the theory must be both conceptually and mathematically correct. This is a result that is difficult to accomplish. In most cases it is practically mandatory to approach the conceptual and mathematical issues separately, so that this very complex problem is reduced to more manageable dimensions. We either develop a mathematically correct theory that is conceptually imperfect (a “model”), and then attack the problem of reconciling this theory with the conceptual aspects of the phenomena in question, or alternatively, develop a theory that is conceptually correct, as far as it goes, but mathematically imperfect, and then attack the problem of accounting for the mathematical forms and magnitudes of the physical relations.

As matters now stand in conventional science, the requirement of conceptual validity is by far the most difficult to meet. With the benefit of the mathematical techniques now available it is almost always possible to devise an accurate, or nearly accurate, mathematical representation of a physical relation on the basis of those physical factors that are *known* to enter into the particular situation, and the *currently accepted concepts* of the nature of these factors. The prevailing policy, therefore, is to give priority to the mathematical aspects of the phenomena under consideration. Vigorous mathematical analysis is applied to models which admittedly represent only certain portions of the phenomena to which they apply, and which, as a consequence, are conceptually incorrect, or at least incomplete. Attempts are then made to modify the models in such a way that they move closer to conceptual validity while maintaining their mathematical validity.

There is a sound reason for following this “mathematics first” policy in the normal course of physical investigation. The initial objective is usually to arrive at a result that is useful in practical application; that is, something that will produce the correct mathematical answers to practical problems. From this standpoint, the issue of conceptual validity is essentially irrelevant. However, scientific investigation does not end at this point. Our inquiry into the subject matter is not

complete until we (1) arrive at a conceptual understanding of the physical phenomena under consideration, and (2) establish the nature of the relations between these and other physical phenomena.

A mathematical relation that is unexplained conceptually is of little or no value toward accomplishing these objectives. It cannot be extrapolated beyond the range for which its validity has been experimentally or observationally verified without running the risk of exceeding the limits of its applicability (as will be demonstrated in Volume III). Nor can it be extended to any area other than the one in which it originated. As it happens, however, many physical problems have resisted all attempts to discover the conceptually correct explanations. Many of the frustrated theorists have reacted by abandoning the effort to achieve conceptual validity, and are now contending that mathematical agreement between theory and observation constitutes “experimental verification.” Obviously this is not true. Such a “verification,” or any number of similar mathematical correlations, tell us only that the theory is *mathematically* correct. As has been emphasized at several points in the preceding discussion, mathematical validity does not, in any way, assure conceptual validity. It gives no indication whether the interpretation that is being given to the mathematical relations is right or wrong. The inevitable result of the currently prevailing policy is to overload physical science with theories that are mathematically correct but conceptually wrong.

Solutions for the many long-standing problems of physical science clearly cannot be obtained as long as the attacks on the problems are terminated when mathematical agreement is achieved. But even if this defect in present practice is corrected, it is doubtful whether the answers to most of these difficult problems can be obtained by the prevailing method of devising a mathematical solution first, and then looking for a conceptual explanation. The reason is that a valid mathematical expression can be constructed to fit almost any model. As Einstein states the case:

It is often, perhaps even always, possible to adhere to a general theoretical foundation by securing the adaptation of the theory to the facts by means of artificial additional assumptions.³⁹

Consequently, the mathematical expressions cannot be relied upon to furnish the necessary clues to a conceptual understanding.

The important contribution of the Reciprocal System of theory to the solution of these problems is that it enables attacking them from the opposite direction; that is, first arriving at a conceptual understanding by deduction from very general basic relations, and then developing the mathematical aspects of the established conceptual relationships. In other words, instead of getting a mathematical answer and then looking for a conceptual explanation to fit it, we start by getting a conceptual answer and then look for a mathematical way of expressing it. In general, this is a much simpler procedure, but it could not be utilized on any extensive scale until a unified general theory was available, so that conceptual answers could be obtained by deductive processes. The Reciprocal System of theory satisfies this requirement.

The clarification of the basic aspects of electricity and magnetism provides a dramatic example of the power of this new method of approach to physical problems. It is no longer necessary to deny the existence of answers to the questions listed at the beginning of this chapter, or to content ourselves with pseudo-answers such as the “curved space” explanation of gravitation. Two of these questions, “What is mass?,” and “What is an electric current?,” have already been answered in the previous pages of this and the preceding volume. Those involving magnetism will be answered in the general discussion of that subject which begins with Chapter 19, and the process of induction of charges will be explained in Chapter 18. The answers to all of the other questions in the list will be developed in this present chapter. When these presentations are complete, we will have provided simple and logical explanations for every one of these items with which present-day science is having so much difficulty.

In the universe of motion all physical entities and phenomena are motions, combinations of motions, or relations between motions. It follows that the development of the structure of the theory that describes this universe is primarily a matter of determining just what motions and combinations of motions can exist under the conditions specified in the postulates. Thus far in our discussion of electrical phenomena we have been dealing only with *translational* motion, the movement of electrons through matter, and the various effects of that motion, the mechanical aspects of electricity, so to speak. We will now turn our attention to the electrical phenomena that involve *rotational* motion.

As we saw in Volume I, gravitation is a three-dimensional rotationally distributed scalar motion. Objects having only one or two *effective* dimensions of scalar rotation were found to exist, but these objects, sub-atomic particles, have only a limited role in physical phenomena. In view of the general pattern of generating motions of greater complexity by combining motions of different kinds, the possibility of superimposing one-dimensional or two-dimensional scalar rotation on gravitating objects to produce phenomena of a more complex nature naturally suggests itself. On analyzing the situation, however, we find that the addition of ordinary rotationally distributed motion in less than three dimensions to the gravitational motion would merely modify the magnitude of that motion, and would not result in any new kinds of phenomena.

There is, however, a modification of the rotational distribution pattern that we have not yet explored. Three general types of simple motion (scalar motion of physical locations) have thus far been considered: (1) translation, (2) linear vibration, and (3) rotation. We now need to recognize that there is a fourth type: *rotational vibration*, a motion that is related to rotation in the same way that linear vibration is related to translational motion. Vectorial motion of this type is uncommon—the motion of the hairspring of a watch is an example—and it is largely ignored in conventional physical thought, but it plays an important part in the basic motion of the universe.

At the atomic level, rotational vibration is a rotationally distributed scalar motion that is undergoing a continuous change from outward to inward and vice versa. As in linear vibration, the change of scalar direction must be continuous and uniform in order to be permanent. Like the motion of the photon of radiation, it is therefore a simple harmonic motion. As noted in the discussion of thermal motion in Chapter 5, when such a simple harmonic motion is added to an existing motion it is coincident with that motion (and therefore ineffective) in one of the scalar directions, and has an effective magnitude in the other scalar direction. Every added motion must conform to the rules for the combination of scalar motions that were set forth in Volume I. On this basis, the effective scalar direction of a self-sustaining rotational vibration must be outward, in opposition to the inward rotational motion with which it is associated. A similar

addition with an inward scalar direction is not stable, but can be maintained by an external influence, as we will see later.

A scalar motion in the form of a rotational vibration will now be identified as a *charge*. A one-dimensional motion of this type is an *electric charge*. In the universe of motion, any basic physical phenomenon such as a charge is necessarily a motion, and the only question to be answered by an examination of its place in the physical picture is what kind of a motion it is. We find that the observed electric charge has the properties that the theoretical development identifies as those of a one-dimensional rotational vibration, and we can therefore equate the two.

It is interesting to note that conventional science, which has been at so much of a loss to explain the origin and nature of the charge, does recognize that it is scalar. For instance, W. J. Duffin reports that experiments which he describes show that “charge can be specified by a single number,” thus justifying the conclusion that “charge is a *scalar* quantity.”⁴⁰

However, in current physical thinking this electric charge is regarded as one of the fundamental physical entities, and its identification as a motion will no doubt be a surprise to many persons. It should therefore be emphasized that this is not a peculiarity of the theory of the universe of motion. Irrespective of our findings, based on that theory, a charge is necessarily a motion on the basis of the *definitions* that are employed in conventional physics, a fact that is disregarded because it is inconsistent with present-day theory. The key factor in this situation is the definition of *force*. It was brought out in Chapter 12 that force is a property of motion, not something of a fundamental nature that exists in its own right. An understanding of this point is essential to the development of the theory of charges, and some further consideration of the relevant facts is therefore appropriate in the present connection.

For application in physics, force is defined by Newton’s second law of motion. It is the product of mass and acceleration, $F=ma$. Motion, the relation of space to time, is measured on an individual mass unit basis as speed, or velocity, v , (that is, *each* unit moves at this rate), or on a collective basis as momentum, the product of mass

and velocity, mv , formerly called by the more descriptive name “quantity of motion.” The time rate of change of the magnitude of the motion is dv/dt (acceleration, a) in the case of the individual unit, and $m dv/dt$ (force, ma) when measured collectively. Thus force is, in effect, defined as the time rate of change of the magnitude of the total quantity of motion, the “quantity of acceleration” we might call it. From this definition it follows that a force is *a property of a motion*. It has the same standing as any other property, and is not something that can exist as an autonomous entity.

The so-called “fundamental forces of nature,” the presumably autonomous forces that are currently being called upon to explain the origin of the basic physical phenomena, are necessarily properties of underlying motions; they *cannot* exist as independent entities. Every “fundamental force” must originate from a fundamental motion. This is a logical requirement of the definition of force, and it is true regardless of the physical theory in whose context the situation is viewed.

Present-day physical science is unable to identify the motions that the definition of force requires. An electric charge, for instance, produces an electric force, but so far as can be determined from observation, it does this on its own initiative. There is no indication of any antecedent motion. This apparent contradiction of the definition of force is currently being handled by ignoring the requirements of the definition, and treating the electric force as an entity generated in some unspecified way by the charge. The need for an evasion of this kind is now eliminated by the identification of the charge as a rotational vibration. It is now clear that the reason for the lack of any evidence of a motion being involved in the origin of the electric force is that *the charge itself is the motion*.

An electric charge is thus a one-dimensional analog of the three-dimensional motion of an atom or particle that we identified as mass. The space-time dimensions of mass are $t^{3/3}$. In one dimension this is $t/3$. Rotational vibration is a motion similar to the rotation that constitutes mass, differing only in the periodic reversal of scalar direction. It follows that the electric charge, a one-dimensional rotational vibration, also has the dimensions $t/3$. The dimensions of

the other electrostatic quantities can be derived from those of charge. The *electric field intensity*, a quantity that plays an important part in many of the relations involving electric charges, is the charge per unit area, $t_s \times 1/s^2 = t_s^3$. The product of field intensity and distance, $t_s^3 \times s = t_s^2$, is a force, the *electric potential*.

For the same reasons that apply to the production of a gravitational field by a mass, the electric charge is surrounded by a force field. However, there is no interaction between mass and charge. As brought out in Chapter 12, a scalar motion that alters the separation between A and B can be represented in the reference system either as a motion AB (a motion of A toward B) or a motion BA (a motion of B toward A). Thus the motions AB and BA are not two separate motions; they are merely two different ways of representing the *same* motion in the reference system. This means that scalar motion is a mutual process, and cannot take place unless the objects A and B are capable of the *same kind* of motion. Consequently, charges (one-dimensional motions) interact only with charges, masses (three-dimensional motions) only with masses.

The linear motion of the electric charge analogous to gravitation is subject to the same considerations as the gravitational motion. As noted earlier, however, it is directed outward rather than inward, and therefore cannot be added directly to the basic vibrational motion in the manner of the rotational motion combinations. This restriction on outward motion is due to the fact that the outward progression of the natural reference system, which is always present, extends to the full unit of outward speed, the limiting value. Further outward motion can be added only after an inward component has been introduced into the motion combination. A charge can therefore exist only as an addition to an atom or sub-atomic particle.

Although the scalar direction of the rotational vibration that constitutes the charge is always outward, both positive (time) displacement and negative (space) displacement are possible, as the rotational speed may be either above or below unity, and the rotational vibration must oppose the rotation. This introduces a rather awkward question of terminology. From a logical standpoint, a rotational vibration with a space displacement should be called a negative charge, since it

opposes a positive rotation, while a rotational vibration with a time displacement should be called a positive charge. On this basis, the term “positive” would always refer to a time displacement (low speed), and the term “negative” would always refer to space displacement (high speed). Use of the terms in this manner would have some advantages, but so far as the present work is concerned, it does not seem advisable to run the risk of adding further confusion to explanations that are already somewhat handicapped by the unavoidable use of unfamiliar terminology to express relationships not previously recognized. For present purposes, therefore, current usage will be followed, and the charges on positive elements will be designated as positive. This means that the significance of the terms “positive” and “negative” with respect to rotation is reversed in application to charge.

In ordinary practice this should not introduce any major difficulties. In this present discussion, however, a definite identification of the properties of the different motions entering into the combinations that are being examined is essential for clarity. To avoid the possibility of confusion, the terms “positive” and “negative” will be accompanied by asterisks when used in the reverse manner. On this basis, an electropositive element, which has low speed rotation in all scalar dimensions, takes a positive* charge, a high speed rotational vibration. An electronegative element, which has both high speed and low speed rotational components, can take either type of charge. Normally, however, the negative* charge is restricted to the most negative elements of this class, those of Division IV.

Many of the problems that arise when scalar motion is viewed in the context of a fixed spatial reference system result from the fact that the reference system has a property, *location*, that the scalar motion does not have. Other problems originate for the inverse reason: scalar motion has a property that the reference system does not have. This is the property that we have called *scalar direction*, inward or outward.

We can resolve this latter problem by introducing the concept of positive and negative reference points. As we saw earlier, assignment of a reference point is essential for the representation of a scalar motion in the reference system. This reference point then constitutes the zero point for measurement of the motion. It will be either a

positive or a negative reference point, depending on the nature of the motion. The photon originates at a negative reference point and moves outward toward more positive values. The gravitational motion originates at a positive reference point and proceeds inward toward more negative values. If both of these motions originate at the same location *in the reference system*, the representation of both motions takes the same form in that system. For example, if an object is falling toward the earth, the initial location of that object is a positive reference point for purposes of the gravitational motion, and the scalar direction of the movement of the object is inward. On the other hand, the reference point for the motion of a photon that is emitted from that object and moves along exactly the same path in the reference system is negative, and the scalar direction of the movement is outward.

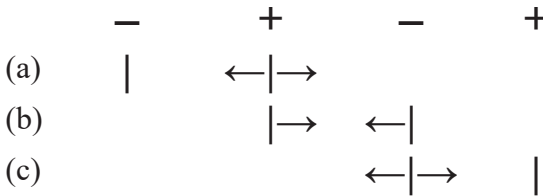
One of the deficiencies of the reference system is that it is unable to distinguish between these two situations. What we are doing in using positive and negative reference points is compensating for this deficiency by the use of an auxiliary device. This is not a novel expedient; it is common practice. Rotational motion, for instance, is represented in the spatial reference system with the aid of an auxiliary quantity, the number of revolutions. Ordinary vibrational motion can be accurately defined only by a similar expedient. Scalar motion is not unique in its need for such auxiliary quantities or directions; in this respect it differs from vectorial motion only in that it has a broader scope, and therefore transcends the limits of the reference system in more ways.

Although the scalar direction of the rotational vibration that constitutes the electric charge is always outward, positive* and negative* charges have different reference points. The motion of a positive* charge is outward from a positive reference point toward more negative values, while that of a negative* charge is outward from a negative reference point toward more positive values. Thus, as indicated in the accompanying diagram, Fig. 20, while two positive* charges (line a) move outward from the same reference point, and therefore away from each other, and two negative* charges (line c) do likewise, a positive* charge moving outward from a positive reference point, as in line b, is moving *toward* a negative* charge

that is moving outward from a negative reference point. It follows that like charges repel each other, while unlike charges attract.

As the diagram indicates, the extent of the inward motion of unlike charges is limited by the fact that it eventually leads to contact. The outward movement of like charges can continue indefinitely, but it is subject to the inverse square law, and is therefore reduced to negligible levels within a relatively short distance.

FIGURE 20



Electric charges do not participate in the basic motions of atoms or particles, but they are easily produced in almost any kind of matter, and can be detached from that matter with equal ease. In a low temperature environment, such as that on the surface of the earth, the electric charge plays the part of a temporary appendage to the relatively permanent rotating systems of motions. This does not mean that the role of the charges is unimportant. Actually, the charges often have a greater influence on the outcome of physical events than the basic motions of the atoms of matter that are involved in the action, but from a structural standpoint it should be recognized that the charges come and go in much the same manner as the translational (kinetic or thermal) motions of the atom. Indeed, as we will see shortly, charges and thermal motion are, to a considerable degree, interconvertible.

The simplest type of charged particle is produced by imparting one unit of one-dimensional rotational vibration to the electron or positron, which have only one unbalanced unit of one-dimensional rotational displacement. Since the effective rotation of the electron is negative, it takes a negative* charge. As indicated in the description of the sub-atomic particles in Volume I, each uncharged electron has two vacant dimensions; that is, scalar dimensions in which there is no effective rotation. We also saw earlier that the basic units of matter,

atoms and particles, are able to orient themselves in accordance with their environments; that is, they assume the orientations that are compatible with the effective forces in those environments. When produced in free space—as, for instance, from the cosmic rays—the electron avoids the restrictions imposed by its spatial displacement (such as the inability to move through space) by orienting in such a way that one of its vacant dimensions coincides with the dimension of the reference system. It can then occupy a fixed position in the natural system of reference indefinitely. In the context of a stationary spatial reference system, this uncharged electron, like the photon, is carried outward at the speed of light by the progression of the natural reference system.

If this electron enters a new environment, and becomes subject to a new set of forces, it can reorient itself to conform to the new situation. On entry into a conducting material, for instance, it encounters an environment in which it is able to move freely, inasmuch as the speed displacement in the motion combinations that constitute matter is predominantly in time, and the relation of the space displacement of the electron to the atomic time displacement is motion. Furthermore, the environmental factors favor such a reorientation; that is, they favor an increase in speed above the previous unit level in a high speed environment, and a decrease in a low speed environment. The electron therefore reorients with its active displacement in the dimension of the reference system. This is either a spatial or a temporal reference system, depending on whether the speed is below or above unity, but the two systems are effectively parallel. They are actually two sections of a single system, as they represent the same one-dimensional motion in two different speed ranges.

Where the speed is above unity, the representation of the variable magnitude is in the temporal coordinate system, and the fixed position in the natural reference system appears in the spatial coordinate system as a movement of the electrons (the electric current) at the speed of light. Where the speed is below unity, these representations are reversed. It does not follow that the progression of the electrons along the conductor takes place at these speeds. In this respect, the electron aggregate is similar to a gas. The individual electrons are moving at high speeds, but in random directions. Only the net excess

of the motion in the direction of the current flow—the electron drift, as it is usually called—is effective as a unidirectional movement.

This idea of an “electron gas” is generally accepted in present-day physics, but it is conceded that “The simple theory runs into greater difficulties when examined in more detail.”⁴¹ As noted previously, the prevailing assumption that the electrons of this electron gas are derived from the structures of the atoms encounters many problems. There is also a direct contradiction in the specific heat values. “The electron gas would be expected to contribute an extra $\frac{3}{2}R$ per mole to the specific heat of metals,”⁴¹ but no such specific heat increment is found experimentally.

The theory of the universe of motion supplies the answers to both of these problems. The electrons whose movement constitutes the electric current are not derived from the atoms, and are not subject to the restrictions that apply to this origin. The answer to the specific heat problem is provided by the nature of the electron motion. The motion of these uncharged electrons (units of space) through the matter of the conductor is equivalent to motion of the matter through extension space. At a given temperature, the atoms of matter have a certain speed relative to space. It is immaterial whether this is extension space or electron space. The motion through electron space (movement of the electrons) is part of the thermal motion, and the specific heat due to this motion is part of the specific heat of the atom, not something separate.

Once the reorientation of the electrons takes place in response to the environmental factors, it cannot be reversed against the forces associated with those factors. The electrons therefore cannot leave the conductor in the uncharged state. The only active property of an uncharged electron is a space displacement, and the relation of this space to extension space is not motion. However, an electron can escape from the conductor by acquiring a charge. A combination of rotational motions (an atom or particle) with a net displacement in space (a speed greater than unity) can move only in time, as indicated earlier, and one with a net displacement in time (a speed less than unity) can move only in space, as motion is a relation *between* space and time. But unit speed (the natural zero or datum level) is

unity in *both* time and space. It follows that a motion combination with a net speed displacement of zero can move in *either* time or space. Acquisition of a unit negative* charge (actually positive in character) by the electron, which, in its uncharged state has a unit negative displacement, reduces the net speed displacement to zero and allows the electron to move freely in either space or time.

The production of a charged electron in a conductor requires only the transfer of sufficient energy to an uncharged electron to bring the existing kinetic energy of that particle up to the equivalent of a unit charge. If the electron is to be projected into space, an additional amount of energy is required to break away from the solid or liquid surface and to overcome the pressure exerted by the surrounding gas. At energies below this level, the charged electrons are confined to the conductor in the same manner as when they are uncharged.

The necessary energy for the production of the charge and the escape from the conductor can be supplied in a number of ways, each of which therefore constitutes a method of producing freely moving charged electrons. A convenient and widely used method furnishes the required energy by means of a voltage differential. This increases the translational energy of the electrons until they meet the requirements. In many applications the necessary increment of energy is minimized by projecting the newly charged electrons into a vacuum rather than requiring them to overcome a gas pressure. The cathode rays used in x-ray production are streams of charged electrons projected into a vacuum. The use of a vacuum is also a feature of the thermionic production of charged electrons, in which the necessary energy is imparted to the uncharged electrons by means of heat. In photoelectric production, the energy is absorbed from radiation.

Existence of the electron as a free charged unit is usually of brief duration. Within a short time after it has been produced by one transfer of energy and ejected into space, it again encounters matter and enters into another energy transfer by means of which the charge is converted back into thermal energy or radiation, and the electron reverts to the uncharged condition. In the immediate neighborhood of an agency that is producing charged electrons, both

the creation of charges and the reverse process that transforms them back into other types of energy are going on simultaneously. One of the principal reasons for the use of a vacuum in electron production is to minimize the loss of charges by way of this reverse process.

Charged electrons in space can be observed—that is, detected by various means—and because of the presence of the charges they are subject to electric forces. This enables control of their motions, and unlike its elusive uncharged counterpart, the charged electron is an observable entity that can be manipulated to produce physical effects of various kinds.

It is not feasible to isolate and examine individual charged electrons in matter as we do in space, but we can recognize the presence of such particles by evidence of freely moving charges within the material aggregates. Aside from the special characteristics of charges, these charged electrons in matter have the same properties as the uncharged electrons. They travel readily through good conductors, less readily through poor conductors, they move in response to voltage differences, they are restrained by insulators, which are substances that do not have the necessary open dimensions to allow free electron motion, and so on. In their activities in and around aggregates of matter, these charged electrons are known as *static electricity*.

CHAPTER 14

The Basic Forces

As brought out in the preceding chapter, the development of a purely deductive theory of the physical universe has enabled reversing the customary procedure in scientific investigation. Instead of deriving mathematical relations applicable to the phenomena under consideration, and then looking for an explanation of the mathematics, we are now able, by deduction from very general premises, to derive a theory that is conceptually correct, and then look for an accurate mathematical representation of the theory. This is a much more efficient procedure, for reasons that were previously explained, but it does not necessarily follow that completing the task by solving the mathematical problems will be free from difficulty. In some cases, the search for the correct mathematical statement will require expenditure of a great deal of time and effort. During the course of this extended investigation there will be some defects in the mathematical “models” that are being used, just as there are defects in the conceptual models that are utilized in current practice.

The original development of the theory of the universe of motion, prior to the first publication in 1959, answered a number of the physical questions with which conventional physical science was (and still is) unable to deal. Atoms and sub-atomic particles were identified as combinations of scalar motions, and gravitation was identified as the inward translational manifestation of these motions. Electric charges were identified as one-dimensional motions of an oscillating character superimposed on the basic motion combinations, with similar translational (scalar) resultants. The basic forces were identified as the force aspects of these basic motions.

These identifications answered the questions as to how the forces are produced and the nature of the originating entities. They also answered the problem of explaining how the gravitational and electrostatic effects

are (apparently) transmitted, and accounting for the instantaneous nature of the apparent transmission. Like many of the other answers to long-standing problems that have emerged from the development of this theory, the answer to the transmission problem took an unexpected form. From the postulates of the theory we deduce that each mass and charge follows its own course, and the apparent transmission is merely a result of the fact that the motion is scalar, and therefore has either an inward scalar direction, carrying all objects of these classes toward each other, or an outward scalar direction, carrying all such objects away from each other. There is no transmission of the effects, and hence no transmission time is involved.

As might be expected, the answers to most of these problems were initially incomplete, and the history of the theory since 1959 has been that of a progressive increase in understanding in all physical areas, during which one after another of the remaining issues has been clarified. In some cases, such as the atomic rotation, the mathematical aspects of the problems presented no particular difficulty, and the points at issue were conceptual. In other instances, the difficulties were primarily concerned with accounting for the mathematical forms of the theoretical relations and their numerical values.

The most troublesome problem of the latter kind has been that of the force equations. The force between electric charges can be calculated by means of the *Coulomb equation*, $F = QQ'/d^2$, which states that, when expressed in appropriate units, the force is equal to the product of the (apparently) interacting charges divided by the square of the distance between them. Aside from the numerical coefficients, this Coulomb equation is identical with the equation for the gravitational force that was previously discussed, and, as we will see later, with the equation for the magnetostatic force as well.

Unfortunately, these force relations occupy a dead end from a theoretical standpoint. Most basic physical relations have the status of points of departure from which more or less elaborate systems of consequences can be built up step by step. Correlation of these consequences with each other and with experience then serves either to validate the theoretical conclusions or to identify whatever errors or inadequacies may exist. No such networks of connections have been

identified for these force equations, and this significant investigative aid has not been available to those who have approached the subjects theoretically. The lack of an explanation has not been as conspicuous in the case of the electric force, as the Coulomb equation, which expresses the magnitude of this force, is stated in terms of quantities derived from the equation itself, but there is an embarrassing lack of theoretical understanding of the basis for the relation that is expressed mathematically in the gravitational equation. Without such an understanding the physicists have been unable to tie this equation into the general structure of physical theory. As expressed in one physics textbook, "Newton's law of universal gravitation is not a defining equation, and cannot be derived from defining equations. It represents an *observed relationship*."

The problems involved in application of the theory of the universe of motion to the gravitational relations were no less formidable, and the early results of this application were far from satisfactory. Ordinarily, the results of incomplete investigations of this kind have not been included in the published material. The opportunities for publication of the findings of this investigation have been severely limited, and the material released for publication has therefore been confined, in general, to those results that have been established as correct both mathematically and conceptually, within the limits to which the investigation has been carried. If the gravitational motion and force were matters of an ordinary degree of importance, the best policy probably would have been to put the unsatisfactory results aside for the time being, and to wait for further developments in related areas to clarify the general situation enough to make further progress in the gravitational area possible. But because of the fundamental nature of the gravitational relations it has been necessary to make extensive use of them, in whatever form they might happen to be, as the theoretical investigation progressed. The previous publications, including the first volume of this present series, have therefore contained some of the tentative, and only partially correct, results of the earlier studies. However, much new light is being thrown on the subject matter by the continuing advances that are being made in related areas, and the status of the gravitational theory is consequently being updated in each new publication.

At the time of the first studies, the most obvious need was a clarification of the dimensions of the equations. Overall dimensional consistency is something that has never been attained by conventional physics. In some areas, such as mechanics, the currently recognized relations are dimensionally consistent, but in many other areas the dimensional confusion is so widespread that it has led to the previously mentioned conclusion that a rational system of dimensions for all physical quantities is impossible.

The present standard practice is to cover up the discrepancies by assigning dimensions to the numerical constants in the equations. Thus the gravitational constant is asserted to have the dimensions $\text{dyne-cm}^2/\text{gram}^2$. Obviously, this expedient is illegitimate. Whatever dimensions enter into physical expressions are properties of the physical entities that are involved, not properties of numbers. Dimensions are excluded from numbers by definition. Wherever, as in the gravitational case, an equation cannot be balanced without assigning dimensions to a numerical constant, this is *prima facie* evidence that there is something wrong in the understanding on which the dimensional assignments are based. Either the dimensions assigned to the physical quantities in the equation are incorrect, or the so-called “numerical constant” is actually the magnitude of an unrecognized physical property. Both types of dimensional errors have been encountered in our examination of current thought in the areas covered by our investigation.

One of the powerful analytical tools made available by the theory of the universe of motion is the ability to reduce all physical quantities to terms of space and time only. In order to be correct, an equation must have a space-time balance; that is, both sides of the equation must reduce to the same space-time expression. Another useful analytical tool derived from this theory is the principle of equivalence of units. This principle asserts that, inasmuch as the basic quantities, in all cases, are units of motion, there are no inherent numerical constants in the mathematical equations that represent physical relations, other than what we may call structural constants—values that have definite physical meanings, as, for instance, the number of active dimensions in one of the participating quantities. It follows that if the quantities involved in a valid physical equation are all

expressed in natural units, or the equivalent in the units of another measurement system, the equation is in balance numerically, and no numerical constant is required.

The gravitational equation, in its usual form, fails by a wide margin to meet the test of dimensional consistency, but the general nature of the modifications that have to be made in the dimensional assignments was identified quite early in the investigation. The 1959 publication dealt with this dimensional problem, pointing out the need to reduce the distance term and one of the mass terms to a dimensionless status; that is, to recognize that they are merely ratios. It also emphasized the fact that an acceleration term must be introduced into the equation for dimensional consistency, and showed that this term represents the inherent acceleration of gravitating objects, which is unity, and therefore not perceptible in empirical measurements. Application of the principle of equivalence of natural units was attempted, without much success, but the tentative results of this study included a derivation of the gravitational constant.

By the time *Nothing But Motion* was published twenty years later, the lack of a fully satisfactory interpretation of the gravitational equation had become somewhat embarrassing. Furthermore, the validity of the original derivation of the gravitational constant was challenged by some of the author's associates, and the evidence in its favor was not sufficient to meet that challenge effectively. It was therefore decided to abandon that interpretation, and to look for a new explanation to take its place. In retrospect it will have to be admitted that this 1979 revision was not a well-conceived attack on the problem. It was essentially an attempt to find a mathematical (or at least numerical) solution where the logical development of the theory had met an obstacle. This is the same policy that, as pointed out in Chapter 13, has brought conventional theory up against so many blank walls, and it has turned out to be equally unproductive in the present case. It became increasingly evident that some further study was necessary.

This brings up an issue that has been the subject of some comment. It is our contention that the many thousands of correlations between the observations and the consequences of the postulates of the theory of the universe of motion have established that this theory is a true and accurate representation of the actual physical universe. The

skeptics then want to know how we can arrive at wrong conclusions in some cases, if we are applying a correct theory; why a conclusion reached in the first volume of a series had to be modified even before the second volume was published. The answer, as explained in many of our previous publications, is that while the theory is capable of producing the right answers, if properly applied, it does not necessarily follow that those who are attempting to apply it properly will always be successful in so doing. As stated earlier, an attempt has been made to confine the published material to firmly established items, aside from a few that are specifically identified as somewhat speculative, but nevertheless, some of the conclusions that have been published have subsequently been found to be incomplete, and in a few instances, incorrect.

There is no reason to be apologetic about these few errors and omissions. Present-day physical theory has been in the process of development for centuries, during which a myriad of conclusions that have been reached with respect to details of the theory (or theories) have subsequently had to be abandoned as incorrect. In comparison with this experience, the error rate in the development of the theory of the universe of motion is fantastically low. This is no accident. Inasmuch as all conclusions in all areas are derived deductively from the same set of basic premises, consistency of the interrelations between phenomena, the basic requirement for conceptual validity, is achieved automatically. Those cases in which the developers of the theory are having some trouble merely emphasize the easy and natural way in which solutions for most of the previously unresolved fundamental problems of physical science have emerged from the theoretical development.

The review of the gravitational situation that was recently undertaken was able to take advantage of some very significant advances that have been made in our understanding of the details of the universe of motion—that is, in the consequences of the postulates—in the years that have elapsed since publication of Volume I in 1979. Chief among these is the clarification of the nature and properties of scalar motion, discussed in Chapter 12, and covered in more detail in *The Neglected Facts of Science*. The improvement in understanding of this type of motion has thrown a great deal of new light on the force relations. It

is now clear that the differences between the basic types of forces that were recognized from the start of the investigation as dimensional in nature are differences in the number of *scalar* dimensions involved, rather than geometric dimensions of space. This provides simple explanations for several of the issues that had been matters of concern in the earlier stages of the theoretical development.

The significant conceptual change here is in the nature of the relation between motion and its representation in the reference system. In previous physical thought motion was regarded as a change of position in a specifically defined physical space (Newton) or spacetime (Einstein) during a specific physical time. This physical space and time thus constitute a background, or container. Changes of position due to motion relative to the spatial background are assumed to be capable of representation by vectors (or tensors of higher rank). In the theory of the universe of motion, on the other hand, space and time have physical existence only as the reciprocally related components of motion, and the three-dimensional space of our ordinary experience is merely a reference system, not a physical container. Furthermore, the development of the details of the theory in the preceding pages of this and the earlier volume shows that the spatio-temporal reference system which combines the three-dimensional spatial frame of reference with the time magnitudes registered on a clock, is incapable of representing the full range of existing motions. Some motions cannot be represented in their true character. Others cannot be represented in this reference system at all.

The deficiency of the reference system with which we are particularly concerned at this time is its inability to represent multi-dimensional scalar motion. This inability of the reference system to represent more than one scalar dimension of motion explains why the forces exerted by charges and masses are all one-dimensional, irrespective of the number of scalar dimensions applicable to the inherent motion of the charge or mass. Only one of these scalar dimensions is coincident with the dimension of the reference system, and the motion in this dimension is therefore the only one that can be represented in the reference system. As indicated earlier, this limitation on the capacity of the reference system is the reason for the great disparity

in magnitude between the basic forces. The *total* magnitudes of the electric and gravitational forces are actually the same, but only the motion in the dimension of the reference system is effective. In our gravitationally bound system, the dimensional ratio (in cgs units) is 3×10^{10} . Thus the electric force, which is one-dimensional, and therefore fully effective, is relatively strong. The gravitational force actually has the same total strength, but it is distributed over three scalar dimensions, only one of which coincides with the dimension of the reference system. The *effective* gravitational force is therefore weaker than the effective electrostatic force by the factor 9×10^{20} .

It should be noted, however, that the difference in the number of effective scalar dimensions has this effect on the relative magnitude of the forces only because it is applied to the very large value of the unit of speed, the relation between the sizes of the units in which we measure space and time. This, in turn, is a consequence of our position in a gravitationally bound system that is moving inward in space at a high speed, opposing the spatial component of the outward progression of the natural reference system. The *net* motion of the gravitating system in space is relatively small, while the motion in time proceeds at the full speed of the progression. Thus we experience a small change in space coincidentally with a very large change in time. We assign values to the units of these quantities that reflect the manner in which we experience them, and on this basis we have defined a unit of time (in the cgs system) that is 3×10^{10} times as large as our unit of space. Our unit of speed is then 3×10^{10} space units (centimeters) per unit of time (second).

As can be seen from the foregoing, the magnitude that we assign to the unit of speed, the speed of light, customarily represented by the symbol c , is not an inherent property of the universe (although the magnitude of the speed itself is). The general range within which this value will fall is determined by our position in a system of gravitating objects, and the specific value within these limits is assigned arbitrarily. Any change in the unit of either space or time that is not counterbalanced by an equivalent change in the other alters the value of c , in our measurement system, and the relation between the magnitudes of the electric and gravitational forces, c^2 , is changed accordingly. (The electric force is usually asserted to be

10^{39} or 10^{40} times as strong as the gravitational force, but this figure is based on a set of erroneous assumptions.)

The further clarification of the mutual nature of scalar motion accomplished in the most recent studies has also thrown a very significant additional light on the force situation. As brought out in Chapter 12, it is now evident that a scalar motion AB cannot be distinguished, in the absence of a fixed coupling to the reference system, from a scalar motion BA. This means that in considering the mutual gravitational motion of two masses we are dealing with only one motion, the representation of which in the reference system depends on external factors.

On this basis, the expression mm' in the gravitational equation is not a product of two masses, but the product of one mass and the *number* of units of mass in the interacting object. Likewise, the distance term, s^2 , is a pure number, the ratio of s^2 units to 1^2 unit. Thus the only dimensional quantity that appears in the equation, aside from the resultant force, is one of the mass terms. This result of the current study confirms the original finding reported in the 1959 publication. It likewise confirms the earlier finding that another dimensional term, a unit of acceleration, must be inserted into the equation to produce a dimensional balance. Force in general is the product of mass and acceleration. It follows that the expression for any *particular* force must reduce to $F=ma$ when all dimensions are properly assigned. The existence of the acceleration term is not apparent without a theoretical analysis because the gravitational acceleration is unity, and therefore has no effect on the numerical result.

The difficulties that have previously been experienced in applying the principle of the equivalence of natural units to the gravitational equation are now seen to have been due to an inadequate understanding of the manner in which the dimensionless terms in the equation should be treated when the statement of the unit equivalence is formulated. We now recognize that these terms vanish if they are given unit value in the system of measurement in which the values of the dimensionless terms are stated, unless some structural factor is specifically applicable. However, the use of an arbitrary mass unit in the conventional measurement systems introduces a complication, as it means that two different systems of units are actually being used.

As we saw in the discussion of physical fundamentals in Volume I, all physical quantities, including mass, can be expressed in terms of units of space and time only. It follows that when an arbitrary unit is used for the measurement of mass, we are expressing the mass and the acceleration in different measurement systems. This is equivalent to introducing a numerical factor into whatever physical relations may be involved: the ratio between the sizes of the respective units.

Introduction of this factor does not affect the numerical balance of an equation as long as both sides of the equation contain the same number of mass terms, but in the gravitational equation $F = kmm'/d^2$, there are two mass terms on one side of the equation, while the force, the lone term on the other side, contains only one mass term ($F = ma$). In order to balance the equation numerically, a correction factor must be applied to convert the extra mass term to the units applicable to space and time. The ratio of the natural space-time unit of mass to the arbitrary mass unit is the required correction factor. Together with whatever structural factors are applicable to the equation, it constitutes the *gravitational constant*.

The ratio of the natural unit of mass in the cgs system to the arbitrary unit, the gram, was evaluated in Volume I as 2.236055×10^{-8} . It was also noted in that earlier volume that the factor 3 (evidently representing the number of effective dimensions) enters into the relation between the gravitational constant and the natural unit of mass. The gravitational constant is then $3 \times 2.236055 \times 10^{-8} = 6.708165 \times 10^{-8}$ (with a small adjustment that will be considered shortly).

To apply the principle of equivalence of natural units to the gravitational equation, the dimensionless quantities m' and d^2 are given unit value in terms of the conventional measurement systems, so that they vanish from the equation. The dimensional terms, the mass term m and the acceleration term inserted into the equation, are then stated in the appropriate *natural* units, 1.65979×10^{-24} grams and 1.971473×10^{26} cm/sec², respectively. The natural unit of force derived from these values is 3.27223×10^2 dynes.

The values thus derived exceed the measured gravitational constant and the previously determined value of unit force by the factor 1.00524. Since it is unlikely that there is an error of this magnitude in the

measurements, it seems evident that there is another, quite small, structural factor involved in the gravitational relation. This is not at all surprising, as we have found in the earlier studies in other areas that the primary mass values entering into physical relations are often subject to modification because of secondary mass effects. The ratio of the unit of secondary mass to the unit of primary mass is 1.00639. The remaining uncertainty in the gravitational values is thus within the range of the secondary mass effects, and will probably be accounted for when a comprehensive study of the secondary mass situation is carried out.

A rather ironic result of the new findings with respect to the gravitational constant, as described in the foregoing paragraphs, is that they have taken us back almost to where we were in 1959. The repudiation of the 1959 result in the 1979 publication as a consequence of the criticism levied against it is now seen to have been a mistake. In the light of the additional information now available it appears that the shortcoming of the original results was not that they were wrong, but that they were incomplete and not adequately supported with explanations and confirmatory evidence, and were therefore vulnerable to attack. The more recent work has provided the support that was originally lacking.

Clarification of the gravitational force equation is not only important in itself, but has a further significance in that it opens the door to an understanding of the general nature of all of the primary force equations. Each of these equations is an expression representing the magnitude of the force (apparently) exerted by one originating entity (mass or charge) on another of the same, or equivalent, kind, at a specified distance. All take the same general form as the gravitational equation, $F = kmm'/d^2$.

With the benefit of the information developed in the earlier pages of this chapter, we may now generalize the equation by replacing m with X , which will stand for any distributed scalar motion with the dimensions $(t/s)^n$, and introducing a term Y with the value $1/s \times (s/t)^{n-1}$. The primary force equation is then $F = kXY (X'/d^2)$.

Since only one dimension of an n -dimensional scalar motion is effective in the space of the conventional reference system, the effective

space-time dimensions of the motion participating in the force equation are t_s . By definition, force has the dimensions t_s^2 . The function of the term Y in the primary force equation is to reduce $(t_s)^n$ to t_s and to introduce the term $1/s$ that is necessary to convert t_s to t_s^2 . In the case of the gravitational equation, this involves multiplying by $s^2/t^2 \times 1/s = s/t^2$. These are the dimensions of acceleration. In the Coulomb equation the correction factor Y is merely $1/s$.

The term X'/d^2 is a combination of two ratios, and has unit value in the unit statement of the equation. The numerical constant k is also unity if all quantities are expressed in units that are consistent with the units in which the space and time magnitudes entering into the equation are measured. Where one or more of these quantities are expressed in units of another kind, the difference in the size of the units appears as the value of the numerical constant k . In the gravitational case, for example, the gravitational constant reflects the result of expressing mass in terms of a special unit (grams in the cgs system) rather than in sec^3/cm^3 .

In essence, all that the force equations do is to reduce the scalar motions (mass, charge, etc.) to their effective one-dimensional values, introduce the $1/s$ term that relates the motion to the corresponding force, and correct for any inconsistencies in the units that are employed. It is somewhat of an anticlimax to arrive at such a simple explanation after years of exploring much more complicated hypotheses, but the simplicity of this result is consistent with the general nature of the findings in the basic areas of other physical fields. There are many complex phenomena in nature, to be sure, but throughout the development of the details of the universe of motion we have found that the fundamental relations are quite simple.

As noted earlier, the reference point of a scalar motion, the point in the fixed reference system to which an object in the scalar motion system is coupled, may be in motion vectorially. The mass of this object is a measure of its three-dimensional distributed scalar motion, the inward gravitational motion. The vectorial motion is outward, and in order for it to take place, a portion of the inward gravitational motion must be overcome. The mass is thus also a measure of the magnitude of the resistance to vectorial motion, the *inertia* of the

object. In the light of the points brought out in the preceding pages, it is evident that in these manifestations of mass we are looking at two aspects of the same thing, just as in the case of the rocket, where the quantity of acceleration imparted by the combustion products (the force) is the same as the quantity of acceleration imparted to the rocket.

This point was not recognized by the early investigators because they were not aware of the existence of motion in different scalar dimensions. It appeared to them that two different quantities were involved: a *gravitational mass* and an *inertial mass*. Very accurate measurements showed that these two masses are identical, a finding that the physics of that day could not explain. As one observer says, "Within the framework of classical physics there is no explanation. When attention was directed to the problem, it seemed like a complete mystery."⁴² A step toward a solution of the problem was taken by Einstein. In the absence of an understanding of scalar motion, he was not able to see that gravitation is a motion, but he formulated a "Principle of Equivalence," in which he postulated that gravitation is *equivalent to a motion*. Since he viewed "motion" as synonymous with "vectorial motion," the postulate meant that gravitation is equivalent to an accelerated frame of reference, and it is often expressed in these terms. But such an equivalence is inconsistent with Euclidean geometry. As explained by Tor Gerholm:

If acceleration and gravitation are equivalent, we must apparently also be able to imagine an acceleration field, a field formed by inertial forces. *It is easy to realize that no matter how we try, we will never be able to get such a field* to have the same shape as the gravitational field around the earth and other celestial bodies... If we want to save the equivalence principle... If we want to retain the identity between gravitational and inertial mass, then *we are forced to give up Euclidean geometry!* Only by accepting a non-Euclidean metric will we be able to achieve a complete equivalence between the inertial field and the gravitational field. This is the price we must pay.⁴³

Identification of gravitation as a distributed scalar motion has now thrown an entirely new light on the situation. Gravitation is an accelerated motion, but it is not geometrically equivalent to

an accelerated frame of reference. Einstein's attempt to reconcile these two phenomena by resort to non-Euclidean geometry is misdirected. Whatever mathematical results are obtained by the use of this expedient (actually not very many. As Paul Davies points out, "technical problems of a mathematical nature render all but the simplest systems hopelessly insoluble"⁴⁴) are not indicative of the true relations. The scalar gravitational motion of an object and any vectorial motion that it may possess are quite different in their nature and properties.

In the case of the propagation of radiation, the principal stumbling block for the ether theory was the contradictory nature of the properties that the hypothetical substance "ether" must possess in order to perform the functions that were assigned to it. Einstein's solution was to replace the ether with another entity that was assumed to have no properties, other than an ability to transmit the radiation, an ability which he says we should "take for granted."²⁷

Similarly, the obstacles to accounting for the observed results of the addition of velocities were the existence of absolute magnitudes and fixed spatial coordinate locations. Here, the answer was to deny the reality of absolute magnitudes, and as Einstein says, to "free oneself from the idea that co-ordinates must have an immediate metrical meaning."⁴⁵ Now we find that he deals with the gravitational problem in the same way, loosening the mathematical constraints, rather than looking for a conceptual error. He invents the "equivalent of motion," a hypothetical something which has enough of the properties of motion to enable accounting for the mathematical results of gravitation (at least in principle) without having those properties of vectorial motion that are impossible to reconcile with the observed behavior of gravitating objects. In all of these cases, the development of the theory of the universe of motion has shown that the real reason for the existence of these problems was the lack of some essential information. In the case of the composition of velocities, the missing item was an understanding of motion in time. In the other two cases cited, the problems were consequences of the lack of recognition of the existence of scalar motion.

CHAPTER 15

Electrical Storage

We now turn to a consideration of the *storage* of uncharged electrons (electric current), a subject that was not considered earlier because it was more convenient to wait until after the nature of electric charges was clarified.

The basic requirement for storage is a suitable container. Any conductor is, to some extent, a container. Let us consider an isolated conductor of unit cross section, a wire. This conductor has a length of n units, meaning that it extends through n units of extension space, the space represented in the reference system. Each of these units of the reference system is a location in which a unit of *actual space* (that is, the spatial component of a motion) may exist. In the absence of an externally applied electric voltage, the wire contains a certain concentration of uncharged electrons (actual units of space), the magnitude of which depends on the composition of the material of the conductor, as explained in Chapter 11. If this wire is connected to a source of current, and a very small voltage is applied, more uncharged electrons flow into the wire until all of the units of the spatial reference system that constitute the length of the wire are occupied. Unless the voltage is increased, the inward flow ceases at this point.

When the wire is fully occupied, the aggregate of electrons could be compared to an aggregate of atoms of matter in one of the condensed states. In these states all of the units of extension space within the limits of the aggregate are occupied, and no further spatial capacity is available. But if a pressure is applied, either an internal pressure, as defined in Chapter 4, or an external pressure, the inter-atomic motions are extended into time, and the addition of the spatial equivalent of this time allows more atoms to be introduced into the same section of the extension space represented in the reference system, increasing the density of the matter (the

number of mass units per unit of volume of extension space) beyond the normal equilibrium value.

This ability of physical phenomena to extend into time when further extension into space is prevented is a general property of the universe that results from the reciprocal relation between space and time. The scope of its application is limited, however, to those situations in which a spatial response to an applied force is not possible. In the example just discussed, the compression of solid matter, the obstacle to further inward movement in space is the discrete unit limitation on subdivision. In a wide variety of astronomical phenomena that will be considered in Volume III, the obstacle is the limit on one-dimensional spatial speed. Here, in the electrical storage process, the obstacle is the fixed relation between the unit of actual space and the unit of extension space. An n -unit section of the extension space represented in the reference system can contain n units of actual space, and no more.

If a voltage is applied to force additional electrons into the fully occupied section of wire, the excess electrons are pushed out into time, where they occupy positions in the spatial equivalent of that time. This penetration into time can only be accomplished by application of a force, as the concentration of uncharged electrons in time is already at an equilibrium level. If the voltage is reduced or eliminated, the restoring force tending to bring the electron concentration back into equilibrium reverses the flow, and the excess electrons move back out of the wire. Application of a positive* voltage similarly withdraws electrons from the wire and from equivalent space.

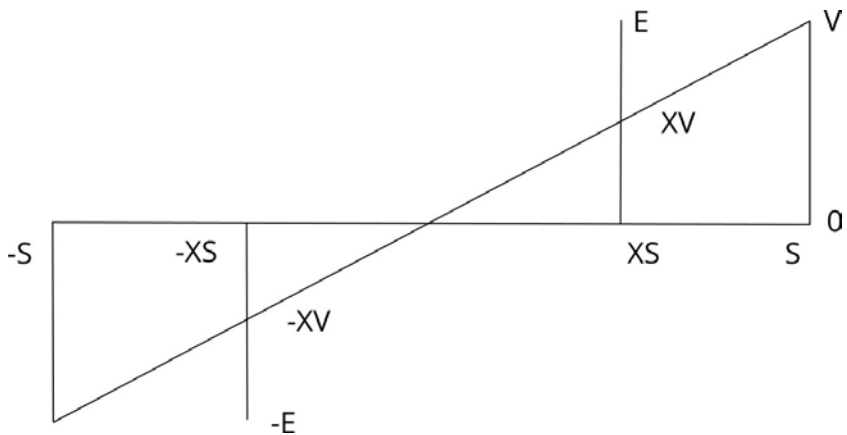
As we have seen in the preceding pages of this and the earlier volume, the region of time beyond the unit of space is two-dimensional. The concentration of excess electrons and the effective voltage therefore decrease in direct proportion to the distance from the wire at a rate determined by basic physical factors and the dimensions of the wire (or other conductor), reaching the zero level at a specific distance.

Let us consider a case in which a conductor is subjected to a voltage differential of 2V, and the voltage in equivalent space surrounding

each terminal reaches zero at a distance s from the terminal. As long as the terminals (electrodes) are separated by a distance greater than $2s$, the electron storage, the quantity of current that can be withdrawn at the positive* terminal and introduced at the negative* terminal, is independent of the location of those terminals. However, if the separation is reduced to less than $2s$, a portion of the volume of equivalent space from which the electrons are being withdrawn coincides with the volume of equivalent space into which electrons are being introduced. The excess and deficiency of electrons in this common volume cancel each other, decreasing the net excess or deficiency at the terminals, and thereby reducing the voltage. This means that where the separation of the terminals is reduced below $2s$, the same amount of storage will take place at a lower voltage, or alternatively, a greater amount of storage will be possible at the same voltage.

The relations involved in the storage of current (uncharged electrons) are illustrated in Fig. 21.

FIGURE 21



When the terminals are separated by the distance $2s$, the full voltage drop, V , takes place at each terminal. The electron excess at the negative* terminal, which we will call E , is proportional to V . If the separation between the terminals is decreased to $2xs$, there is an overlap of the equivalent volumes to which the excess and deficiency of electrons are distributed, as indicated above. The effective voltage then drops to xV . At this point, the electron concentration

corresponding to xV is in the equivalent volume at the negative* terminal, while the balance of the total electron input represented by E is in the common equivalent volume, where the net concentration of excess electrons is zero. If the voltage is reduced, the electrons from the common equivalent volume, and from the volume related to the negative* terminal only, flow out of the system in the same proportions in which they entered. Thus the storage capacity at a separation $2x_s$ and voltage xV is the same as that at a separation $2s$ and voltage V . Generalizing this result, we may say that the storage capacity, at a given voltage, of a combination of positive* and negative* electrodes in close proximity is inversely proportional to the distance between them.

The ability of a conducting wire to accept additional electrons when subjected to a voltage makes it available as a container in which uncharged electrons (units of electric current) can be stored and withdrawn as desired. Such storage has some uses in electrical practice, but it is inconvenient for general use. More efficient storage is made possible by a device that contains the necessary components in a more compact form. In this device, a *capacitor*, two plates, each with an area s^2 , are separated by a distance s' . Each plate is equivalent to s^2 conductors of unit cross section. Thus the storage capacity of a capacitor at a given voltage is directly proportional to the plate area and inversely proportional to the distance between the plates. This storage capacity is called the *capacitance*, symbol C . Since it has the dimensions of space ($s^2/s' = s$), it can be calculated directly from the geometrical dimensions of the capacitor. The centimeter has been use as a unit, although the present practice is to use a special unit, the *farad*.

If a capacitor is connected to a current supply, the effective voltage, a force ($\frac{1}{2}s^2$), pushes the uncharged electrons that constitute the current into the capacitor until the concentration corresponding to that voltage is reached. The space-time dimensions of the product are $\frac{1}{2}s^2 \times s = \frac{1}{2}s$. This is inverse speed, or energy. It is not a charge, on the basis of the definition of charge given in this work, but since electric charge has the dimensions of energy, $\frac{1}{2}s$, the quantity stored is equivalent to charge. To minimize the deviations from currently accepted terminology, we will call it a *capacitor charge*. The magnitude of the storage can be expressed by the equation $Q = CV$, where Q is the capacitor charge, C

is the capacitance, and V is the voltage differential across the plates of the capacitor.

The unit of capacitance, the farad, is defined as one coulomb per volt. The volt is one joule per coulomb. These are units of the SI system, which will be used in most of the subsequent discussion of electricity and magnetism, rather than the cgs system of measurement that is in general use in these volumes, the reason being that a substantial amount of clarification of the physical relations in these areas has been accomplished in very recent years, and most of the current literature relating to these subjects utilizes the SI system.

Unfortunately, this recent clarification of the electrical and magnetic situations has not extended to some of the most fundamental issues, including the many problems introduced into electrical theory by the failure to recognize the existence of uncharged electrons and the consequent lack of distinction between electric quantity and electric charge. As we saw in Chapter 9, the unit of electric quantity is a unit of space (s). We find that the unit of electric charge is a unit of energy ($\frac{1}{2}s$). In current practice, both of these quantities are expressed in the same measurement unit, esu (cgs system) or coulombs (SI system). Now that the electric charge has been introduced into our subject matter, we will have to make the distinction that current theory does not recognize, and instead of dealing only with coulombs, we will have to specify coulombs (s) or coulombs ($\frac{1}{2}s$). In this work the symbol Q , which is currently being used for both quantities, will refer only to electric charge, or capacitor charge, measured in coulombs ($\frac{1}{2}s$). Electric quantity, measured in coulombs (s) will be represented by the symbol q .

Returning now to the question as to the quantities entering into the capacitance, the volt, a unit of force, has the space-time dimensions $\frac{1}{2}s^2$. Since capacitance, as we have now seen, has the dimensions of space, s , the coulomb, as a product of volts and farads has the dimensions $\frac{1}{2}s^2 \times s = \frac{1}{2}s$. But the coulomb as the quotient of joules/volts, has the dimensions $\frac{1}{2}s \times s^2/\frac{1}{2}s^2 = s$. Thus the coulomb that enters into the definition of the farad is not the same coulomb that enters into the definition of the volt. We will have to revise these definitions for our purposes, and say that the farad is one coulomb ($\frac{1}{2}s$) per volt, while the volt is one joule per coulomb (s).

The confusion between quantity (s) and charge (t_s) prevails throughout the electrostatic phenomena. In most cases this does not result in any *numerical errors*, because the calculations deal only with electrons, each of which constitutes one unit of electric quantity, and is capable of taking one unit of charge. Thus the identification of the number of electrons as a number of units of charge instead of a number of units of quantity does not alter the numerical result. However, this substitution does place a roadblock in the way of an understanding of what is actually happening, and many of the relations set forth in the textbooks are incorrect.

For instance, the textbooks tell us that $E = Q/s^2$. E, the electric field intensity, is force per unit distance, and has the space-time dimensions $\text{t}_s^2 \times \text{t}_s = \text{t}_s^3$. The dimensions of Q/s^2 are $\text{t}_s \times \text{t}_s^2 = \text{t}_s^3$. This equation is therefore dimensionally correct. It tells us that, as we would expect, the magnitude of the field is determined by the magnitude of the charge. On the other hand, this same textbook gives the equation expressing the force exerted on a charge by the field as $F = QE$. The space-time dimensions of this equation are $\text{t}_s^2 = \text{t}_s \times \text{t}_s^3$. The equation is therefore invalid. In order to arrive at a dimensional balance, the quantity designated as Q in this equation must have the dimensions of space, so that the equation in space-time form will become $\text{t}_s^2 = s \times \text{t}_s^3$. In this case, then, the Q term is actually q (quantity) rather than Q (charge), and the applicable relation is $F = qE$.

The error due to the use of Q instead of q enters into many of the relations involving capacitance, and has introduced considerable confusion into the theory of these processes. Since we have identified the stored energy, or capacitor charge, as dimensionally equivalent to charge, Q, the capacitance equation in its customary form, $Q = CV$, reduces to $\text{t}_s = s \times \text{t}_s^2$, which is dimensionally consistent. The conventional form of the energy (or work, symbol W) equation is $W = QV$, reflecting the definition of the volt as one joule per coulomb. If CV is substituted for Q in this equation, as would appear to be justified by the relation $Q = CV$, the result is $W = CV^2$. This equation is not dimensionally valid, but it and its derivatives can be found throughout the scientific literature. For instance, the development of theory in this area in one current textbook⁴⁶ begins with the following equation for the potential energy of a charge: $dW = VdQ = Q/CdQ$. As

we saw earlier, the coulomb term in $W = VQ$ is coulombs (s) and the equation is actually $W = Vq$. Thus the textbook equation should read: $dW = V dq = Q/C dq$. The space-time dimensions then are:

$$t/s = t/s^2 \times s = t/s \times 1/s \times s$$

The text then goes on to deduce that “to charge the conductor to a potential $V = Q/C$ we must apply an energy

$$W = \int_0^Q \frac{Q}{C} dQ = \frac{Q^2}{2C} = \frac{1}{2} QV = \frac{1}{2} CV^2 ”$$

Here we again find the lack of distinction between Q and q , followed by a substitution of CV for a Q term that is actually q . With the understanding that CV is equivalent to q , not Q , the correct statement of the equation is:

$$W = \int_0^Q \frac{Q}{C} dq = \frac{Qq}{2C} = \frac{1}{2} qV = \frac{1}{2} (CV)V$$

and the space-time dimensions are:

$$t/s = t/s \times 1/s \times s = t/s \times s \times 1/s = s \times t/s^2 = s \times t/s^2$$

Since the capacitance is area divided by distance, A/s , the text replaces C in the expression CV^2 by A/s , and arrives at the following equalities:

$$\frac{1}{2} CV^2 = \frac{1}{2} \epsilon_0 \frac{A}{s} V^2 = \frac{1}{2} \epsilon_0 \left(\frac{V}{s} \right)^2 As$$

ending with an expression in terms of E , the electric field intensity, and As , the volume occupied by the electric field.

$$= \left(\frac{1}{2} \epsilon_0 E^2 \right) As$$

When rewritten to separate out the expressions that are equivalent to q , the equation becomes:

$$\frac{1}{2} (CV)V = \frac{1}{2} \epsilon_0 \left(\frac{A}{s} \right) V = \frac{1}{2} \epsilon_0 \frac{V}{s} \left(\frac{A}{s} \right) s = \frac{1}{2} \epsilon_0 E \left(\frac{A}{s} \right) s$$

In space-time terms, where (CV) and (A/s) are both equal to q and therefore to s , we obtain:

$$s \times t/s^2 = s \times t/s^2 = t/s^3 \times s \times s = t/s^3 \times s \times s$$

The first column of the accompanying tabulation shows the expressions that are equated to energy in the successive steps in this development. As indicated in the second column, the dimensional error in the first equation carries through the entire sequence, and the space-time dimensions of these expressions remain at t^2/s^2 instead of the correct t/s .

In Textbook		Correct	
QV	$t/s \times t/s^2 = t^2/s^3$	qV	$s \times t/s^2 = t/s$
$Q_C dQ$	$t/s \times 1/s \times t/s = t^2/s^3$	$Q_C dq$	$t/s \times 1/s \times s = t/s$
CV^2	$s \times t^2/s^4 = t^2/s^3$	qV	$s \times t/s^2 = t/s$
$E^2 As$	$t^2/s^6 \times s^2 \times s = t^2/s^3$	$E(q/s^2) As$	$t/s^3 \times s/s^2 \times s^2 \times s = t/s$

The error in this series of expressions was introduced at the start of the theoretical development by a faulty definition of voltage. As indicated earlier, the volt is defined as one joule per coulomb, but because of the lack of distinction between charge and quantity in current practice, it has been assumed that the coulomb entering into this definition is the coulomb of charge, symbol Q. In fact, as brought out in the previous discussion, the coulomb that enters into the energy equation is the coulomb of quantity, which we are denoting by the symbol q. The energy equation, then, is not $W=QV$, but $W=qV$.

The correct terms and dimensions corresponding the those in the first two columns of the tabulation are shown in columns 3 and 4. Here the term Q in the first two expressions, and the term CV which was substituted for Q in the last two have been replaced by the correct term q. As indicated in the tabulation, this brings all four expressions into agreement with the correct space-time dimensions, t/s , of energy. The purely numerical terms in all of these expressions were omitted from the tabulation, as they have no bearing on the dimensional situation.

When the full capacity of the capacitor at the existing voltage is reached, the opposing forces arrive at an equilibrium, and the flow of electrons into the capacitor ceases. Just what happens while the capacitor is filling or discharging is something that the theorists have found very difficult to explain. Maxwell found the concept of a "displacement current" essential for completing his mathematical treatment of magnetism, but he did not regard it as a real current.

“This displacement does not amount to a current,” he says, ” but it is the commencement of a current.” He describes the displacement as “a kind of elastic yielding to the action of the force.”⁴⁷ Present-day theorists find this explanation unacceptable because they have discarded the ether that was fashionable in Maxwell’s day, and consequently have nothing that can “yield” where the plates of a capacitor are separated by a vacuum. The present tendency is to regard the displacement as some kind of a modification of the electromagnetic field, but the nature of the hypothetical modification is vague, as might be expected in view of the lack of any clear understanding of the nature of the field itself. As one textbook puts it, “The displacement current is in some ways the most abstract concept mentioned in this book so far.”⁴⁸ Another author states the case in these words:

If one defines current as a transport of charge, the term displacement current is certainly a misnomer when applied to a vacuum where no charges exist. If, however, current is defined in terms of the magnetic fields it produces, the expression is legitimate.⁴⁹

The problem arises from the fact that while the physical observations and the mathematical analysis indicate that a current is flowing into the space between the plates of the capacitor when that space is a vacuum, as well as when it is occupied by a dielectric, such a current flow is not possible if the entities whose movement constitutes the current are charged electrons, as currently assumed. As stated in the foregoing quotation, there are no charges in a vacuum. This impasse between theory and observation that now prevails is another of the many items of evidence showing that the electric current is *not* a movement of charged particles.

Our analysis shows that the electrons do, in fact, flow into the spatial equivalent of the time interval between the plates of the capacitor, but that these electrons are not charged, and are unobservable in what is called a vacuum. Aside from being only transient, this displacement current is essentially equivalent to any other electric current.

The additional units of space (electrons) forced into the time (equivalent space) interval between the plates increase the total space content. This can be demonstrated experimentally if we introduce a

dielectric liquid between the plates, as the increase in the amount of space decreases the internal pressure, the force per unit area due to the weight of the liquid. For this purpose we may consider a system in which two parallel plates are partially immersed in a tank of oil, and so arranged that the three sections into which the tank is divided by the plates are open to each other only at the bottom of the tank. If we now connect the plates to a battery with an effective voltage, the liquid level rises in the section between the plates. From the foregoing explanation it is evident that the voltage difference has reduced the pressure in the oil. The oil level has then risen to the point where the weight of the oil above the free surface balances the negative increment due to the voltage differential.

Because accepted theory requires the “displacement current” to behave like an electric current without being a current, conventional science has had great difficulty in ascertaining just what the displacement actually is. It is an essential element in Maxwell’s formulations, but some present-day authors regard it as superfluous. “All the physics of dielectrics could be discussed without ever bringing in the displacement vector,”⁵⁰ says Arthur Kip. One of the principal factors contributing to this uncertainty as to its status is that the displacement is customarily defined and treated in electrostatic terms, whereas it is actually a manifestation of current electricity. In Maxwell’s equation for the displacement current, the current density, $1/s^2$, and the time derivative of the displacement, dD/dt , are additive, and are therefore terms of an equivalent nature; that is, they have the same dimensions. The space-time dimensions of current density are $s/t \times 1/s^2 = 1/s_t$. The dimensions of D , the displacement, are then $1/s_t \times t = 1/s$. Its place in the capacitance picture is then evident. In the storage process, units of space, uncharged electrons, are forced into the surrounding equivalent space—that is, the spatial equivalent of time ($t = 1/s$)—and this inverse space, $1/s$, becomes one of the significant quantities with which we must deal.

In the customary electrostatic treatment of the displacement, it is defined as $D = \epsilon_0 E$, where E is the field intensity (an electrostatic concept) and ϵ_0 is the *permittivity* of free space. Since the dimensions of E are $1/s^3$, and we have now found those of D to be $1/s$, the space-time dimensions of permittivity are $1/s \times s^3/t = s^2/t$. In current practice,

however, the permittivity is expressed in farads per meter. This makes it dimensionless, since both the farad and the meter are units of space. In the textbooks the dimensions are given as $M^{-1}L^{-3}T^2Q^2$ (mass, length, time, charge). Analyzing these dimensions into space-time terms shows that they agree with the dimensionless farad per meter unit only if Q^2 is interpreted as Q (coulombs t/s) $\times q$ (coulombs s) in the manner that we found necessary in order to clarify the definition of the farad. But in actual application, the permittivity is usually treated as a quantity with the dimensions s^2/t . This is the result that is obtained if the charge dimensions are interpreted as q^2 . The space-time dimensions are then $s^3/t^3 \times 1/s^3 \times t^2 \times s^2 = s^2/t$.

As an example of the usage of permittivity, one expression derived for the electric field intensity due to a point charge is $E = q/\epsilon_0 s^2$, where s is the distance from the charge. Inasmuch as the electric field intensity has the dimensions t/s^3 , the permittivity must have the dimensions s^2/t in order to balance the equation dimensionally. The correct space-time statement is $t/s^3 = s \times t/s^2 \times 1/s^2$.

As pointed out earlier, numerical coefficients cannot have dimensions in the sense in which this term is used in physics. Dimensions, in this case, are properties of the motions that constitute physical entities and phenomena and they are applicable only to physical quantities. If the permittivity has the dimensions s^2/t , as indicated by Maxwell's equation and by the manner in which it is applied, it is a quantity with a specific physical significance. Like the displacement, it is the inverse of a familiar quantity. Displacement is inverse space, while permittivity is inverse force. The equation $D = \epsilon_0 E$, or $E = \epsilon_0/D$, gives the electric field intensity the dimensions $1/s \times t/s^2 = t/s^3$, the same result we get by dividing potential, t/s^2 by distance, s , in defining this quantity.

We are thus faced with a conflict between the dimensional definition of permittivity expressed in the conventional units and the definition derived from Maxwell's relations, a definition that is consistent with the dimensions of displacement. The relation between the two in space-time terms, which is s^2/t , shows where the difference originates, as this is the relation of the unit of electric current, s , to the electrostatic unit, t/s . The dimensionless farad per meter is an *electrostatic* unit, while the s^2/t dimensions for permittivity relate this

quantity to the *electric current* system. [On this basis capacitance, expressed in farads, has the dimensions s^3/t , since $F = q/\sqrt{V} = s \times s^2/t$. This means that capacitance is the inverse of electric field intensity].^{51a}

Permittivity is of importance mainly in connection with non-conducting substances, or *dielectrics*. If such a substance is inserted between the plates of a capacitor, the capacitance is increased. The rotational motions of all non-conductors contain motion with space displacement. It is the presence of these space components that blocks the translational motion of the uncharged electrons through the time components of the atomic structure, and makes the dielectric substance a non-conductor. Nevertheless, dielectrics, like all other ordinary matter, are predominantly time structures; that is, their net total displacement is in time. This time adds to the time of the reference system, and thus increases the capacitance.

From this explanation of the origin of the increase, it is evident that the magnitude of the increment will vary by reason of differences in the physical nature of the dielectrics, inasmuch as different substances contain different amounts of speed displacement in time, arranged in different geometrical patterns. The ratio of the capacitance with a given dielectric substance between the plates to the capacitance in a vacuum is the *relative permittivity*, or *dielectric constant*, of the substance.

The dielectric constants of most of the common dielectric substances—Class A dielectrics, as they are called—show little variation at low frequencies under ordinary conditions.⁵¹ This indicates that permittivity is an inherent property of the substance, a consequence of its composition and structure, rather than of its relation to the environment. This is consistent with the theoretical explanation given above.

Conventional theories of dielectric phenomena are based on the premise that these phenomena are electrostatic in nature. It should be understood, however, that all theories which depend on the existence of electric charges in electrically neutral materials cannot be other than hypothetical. Furthermore, conventional science has no *comprehensive* electrostatic theory of dielectrics. As expressed by W. J. Duffin:

It is important to realize that calculations of fields due to, and forces on, charge distribution are performed on a *model* and the results compared with experiment...different models are required to account for different sets of experimental results.⁵²

In the model that is applied to the capacitance problem it is assumed (1) that positive* and negative* charges exist in the electrically neutral dielectric, (2) that “small movements of the charges have taken place in opposite directions,” and (3) that these movements produce the “polarization which *we believe* takes place (*italics added*).”⁵³ As this statement by Duffin concedes, there is no direct evidence of the polarization that plays the principal role in the theory. The entire “model” is hypothetical.

Clarification of the dimensions of the quantity known as permittivity eliminates the static charges from the mathematics of the electrical storage process, and thereby cuts the ground out from under all of the electrostatic models. The customary mathematical treatment is carried out in terms of four quantities, the displacement, D , the polarization, P , the electric field intensity, E , and the permittivity, ϵ_0 . These quantities, the investigators tell us, are related by the expression $P = D - \epsilon_0 E$. We have already seen, earlier in this chapter, that the space-time dimensions of D are $1/s$, and those of the permittivity, ϵ_0 , are s^2/λ . The dimensions of the quantity $\epsilon_0 E$ are then $s^2/\lambda \times 1/s^3 = 1/s$. It follows that the dimensions of P are also $1/s$.

We thus find that all of the quantities entering into the dielectric processes are quantities related to the electric current: the electric quantity (s), the capacitance (s), the displacement ($1/s$), the polarization ($1/s$), and the quantity $\epsilon_0 E$, which likewise has the dimensions $1/s$. No quantity with the dimensions of charge ($1/s$) has any place in the mathematical treatment. The *language* is that of electrostatics, using terms such as “polarization,” “displacement,” etc., and an attempt has been made to introduce electrostatic quantities by way of the electric field intensity, E . But it has been necessary to couple E with the permittivity, ϵ_0 , and to use it in the form $\epsilon_0 E$, which, as just pointed out, cancels the electrostatic dimension of E . Electric charges thus play no part in the mathematical treatment.

A similar attempt has been made to bring the electric field intensity, E , into relations that involve the current density. Here, again, the

electrostatic quantity, E , is out of place, and has to be removed mathematically by coupling it with a quantity that converts it into something that has a meaning in electric current phenomena. The quantity that is utilized for this purpose is *conductivity*, symbol σ , space-time dimensions s^2/t^2 . The combination σE has the dimensions $s^2/t^2 \times t/s^3 = 1/st$. These are the dimensions of current density. Like the expression $\epsilon_0 E$, previously discussed, the expression σE has a physical meaning only as a whole. Thus it is indistinguishable from current density. The conventional model brings the field intensity into the theoretical picture, but here, again, it is necessary to remove it by a mathematical device before the theory can be applied in practice.

In those cases where the electric field intensity has been used in dealing with electric current phenomena, *without* introducing an offsetting quantity such as σ or ϵ_0 , the development of theory leads to wrong answers. For example, in their discussion of the “theoretical basis of Ohm’s law,” Bleaney and Bleaney say that “when an electric field strength E acts on a free particle of charge q , the particle is accelerated under the action of the force,” and this “leads to a current increasing at the rate $dJ/dt = n(q^2/m)E$,”⁵⁴ where n is the number of particles per unit volume. The space-time dimensions of this equation are $1/st \times 1/t = 1/s^3 \times s^2 \times s^3/t^3 \times 1/s^3$. Thus the equation is *dimensionally* balanced. But it is *physically* wrong. As the authors admit, the equation “is clearly at variance with the experimental observation.” Their conclusion is that there must be “other forces which prevent the current from increasing indefinitely.”

This fact that the key element of the orthodox theory of the electric current, the hypothesis as to the origin of the motion of the electrons, is “clearly at variance” with the observed facts is a devastating blow to the theory, and not all of its supporters are content to simply ignore the contradiction. Some attempt to find a way out of the dilemma, and produce explanations such as the following:

When a constant electric field E is applied each electron is accelerated during its free path by a force $-eE$, but at each collision it loses its extra energy. The motion of the electrons through the wire is thus a diffusion process and we can associate a mean drift velocity $\bar{v}$ with this motion...⁵⁵

But collisions do not transform accelerated motion into steady flow. If they are elastic, as the collisions of the electrons presumably are, the acceleration in the direction of the voltage gradient is simply transferred to other electrons. If the force Eq actually existed, as present-day electrical theory contends, it would result in accelerating the *average* electron. The authors quoted in reference 54 evidently recognize this point, but they fall back on the prevailing confidence that *something* will intervene to save the “moving charge” theory of the electric current from its multiplicity of problems; there “must be other forces” that take care of the discrepancy. No one wants to face the fact that a direct contradiction of this kind invalidates the theory.

The truth is that this concept of an electrostatic force (Eq) applied to the electron mass is one of the fundamental errors introduced into electrical theory by the assumption that the electric current is a motion of electric charges. As the authors quoted above bring out in the derivation of their electric current equation, such a force would produce an accelerating rate of current flow, conflicting with the observations. In the universe of motion the moving electrons that constitute the electric current are uncharged and massless. The mass that is involved in the current flow is not a property of the electrons, which are merely rotating units of space; it is a property of the matter of the conductor. Instead of an electrostatic force, $\frac{1}{s^2}$, applied to a mass, $\frac{t^3}{s^3}$, producing an acceleration ($F/m = \frac{1}{s^2} \times \frac{s^3}{t^3} = \frac{s}{t^2}$), what actually exists is a mechanical force (voltage, $\frac{1}{s^2}$) applied to a mass per unit time, a resistance, $\frac{t^2}{s^3}$, producing a steady flow, an electric current ($V/R = \frac{1}{s^2} \times \frac{s^3}{t^2} = \frac{s}{t}$).

Furthermore, it is observed that the conductors are electrically neutral even when a current is flowing. The explanation given for this in present-day electrical theory is that the negative* charges which are assumed to exist on the electrons are neutralized by equivalent positive* charges on the atomic nuclei. But if the hypothetical electrostatic charges are neutralized so that no net charge exists, there is no electrostatic force to produce the movement that constitutes the current. Thus, even on the basis of conventional physical theory, there is abundant evidence to show that the moving electrons do not carry charges. The identification of the electric current phenomena with the *mechanical* aspects of electricity that we derive from the

theory of the universe of motion now provides a complete and consistent explanation of these phenomena without recourse to the hypothesis of moving charged electrons.

As noted in Chapter 13, charged electrons are subject to the same forces that apply to their uncharged counterparts, as well as to those specifically appertaining to the charges. It would therefore be theoretically possible to apply a voltage and store these charged electrons in capacitors in the same manner as the uncharged electrons (electric current). In practice, however, the storage of charged electrons is accomplished in a totally different manner. A widely used electrostatic device is the Van de Graaf generator. In this generator charged electrons are produced and sprayed onto a moving belt of insulating material. The belt carries them to a storage unit in the form of a large hollow metal sphere. The electrons pass from the belt into the sphere, gradually building up a potential that may reach a level as high as several million volts.

In our examination of electric current phenomena in the preceding chapters we found that the electrons which constitute the current move from the regions of higher voltage (greater concentration or higher speed of the electrons) to regions of lower voltage. In the Van de Graaf generator, electrons of very low electrostatic potential on the belt pass into a container in which the potential may be in the million volt range. Obviously, we are dealing with two different things, both having the dimensions of force, and both customarily measured in volts, but physically unlike in some important respects.

It should now be evident why the term “potential” was not used in the preceding pages in connection with capacitor storage, or other electric current phenomena. The property of the electric current that we are calling “voltage” is the mechanical force of the current, a force that acts in the same manner as the force responsible for the pressure exerted by a gas. Electrostatic potential, on the other hand, is the radial force of the charges, which decreases rapidly with the distance. The *potential* of a charged electron (in volts) is very large compared to the contribution of the translational motion of that electron to the *voltage*. It follows that even where the potential is in the million volt range, the electron concentration in the storage

sphere, and the corresponding voltage, may be low. In that event, the small buildup of the voltage in the electrode at the end of the belt is enough to push the charged electrons into the storage sphere, regardless of the high electrostatic potential.

Many present-day investigators realize that they cannot account for electric currents by means of electrostatic forces alone. Duffin, for instance, tells us that “In order to produce a steady current there must be, *for at least part of the circuit*, non-electrostatic forces acting on the carriers of charge.”¹³ His recognition that these forces act on “the carriers of charge,” the electrons, rather than on the charges, is particularly significant, as this means that neither the forces nor the objects on which they act are electrostatic. Duffin identifies the non-electrostatic forces as being derived “from electromagnetic induction” or “non-homogeneities, such as boundaries between dissimilar materials, or temperature gradients.”

Since the electric currents available to both the investigators and the general public are produced either by electromagnetic induction or by processes of the second non-electrostatic category mentioned by Duffin (batteries, etc.), the non-electrostatic forces that admittedly must exist are adequate to account for the current phenomena as a whole, and there is no need to introduce the hypothetical electrostatic charge and force. We have already seen that the charge does not enter into the mathematics of the current flow and storage processes. Now we find that it has no place in the qualitative explanation of current flow either.

Addition of these further items of physical and mathematical evidence to those discussed earlier now provides conclusive proof that the *mathematical* structure of theory dealing with the storage of electric current is not a representation of *physical* reality. This is not an isolated case. As pointed out in Chapter 13, the conditions under which scientific investigation is conducted have had the effect of directing the investigation into mathematical channels, and the results that have been attained are almost entirely mathematical. As expressed by Richard Feynman:

Every one of our laws is a purely mathematical statement in rather complex and abstruse mathematics.⁵⁶

The development of this mathematical structure of theory is an outstanding achievement, and it has had very important—even spectacular—practical results. However, these successes have fostered a tendency to forget that mathematics is not physics. It is a useful, perhaps indispensable, tool for the physicist, but physical phenomena are subject to a multitude of limitations that do not apply to the mathematics that are utilized to represent these phenomena, and consequently are not recognized unless they are identified physically. The mathematical representation of space, for example, can be “curved,” or otherwise modified, but this does not, in any way, assure us that physical space can be so modified. That question can be settled only by means of a purely *physical* investigation such as the one reported in this work, which finds that such a modification of extension space is impossible.

CHAPTER 16

Induction of Charges

Clarification of the structure of the gravitational equation and application of the new information to formulation of the primary force equation opens the door to an understanding of the Coulomb equation, $F = QQ'/d^2$, that expresses the electrostatic force. This equation is set up on an equivalent basis without a numerical coefficient; that is, the numerical value of the charge Q is defined by the equation itself. It would seem, therefore, that when the other quantities in the equation, force, F , and distance, d , are expressed in terms of the cgs equivalents of the natural units, Q should likewise take the cgs value of the appropriate natural unit. But the dimensions of charge are $\frac{1}{s}$, and the natural unit of $\frac{1}{s}$ in cgs units is 3.334×10^{-11} sec/cm, whereas the experimental unit of charge has the different numerical value 4.803×10^{-10} . In conventional physics there is no problem here, as the unit of charge is regarded as an independent quantity. But in the context of the theory of the universe of motion, where all physical quantities are expressed in terms of space and time only, it has been a puzzle that we have only recently been able to solve.

One of the new items of information that was derived from the most recent analysis of the gravitational equation, and incorporated into the primary force equation, is that the individual force equations deal with only one force (motion). The force (apparently) exerted by charge A on charge B and the force (apparently) exerted by charge B on charge A are not separate entities, as they appear to be; they are merely different aspects of the same force. The reasons for this conclusion were explained in the gravitational discussion.

A second point, also derived from the gravitational study reported in Chapter 14, although it could have been arrived at independently, is that there is a missing term in the usual statement of each of the force equations. This term, identified as $\frac{1}{s} \times (\frac{s}{t})^{n-1}$ in the primary

force equation, must be supplied in order to balance the equation. In the gravitational equation it is an acceleration term. In the Coulomb equation it is reciprocal space, $\frac{1}{s}$.

Here we encounter a difference between the two equations that we have been examining. In the gravitational equation the unit of mass is defined independently of the equation. In the Coulomb equation, however, the unit of charge is defined by the equation. Consequently, any term that is omitted from the statement of the equation is automatically combined with the charge, instead of having to be introduced separately, as was necessary in the case of the acceleration term of the gravitational equation. The quantity $\frac{1}{s}$, which, as we have just seen, is required for a dimensional balance, therefore becomes a component of the quantity that is called "charge" in the statement of the equation. That quantity is actually $\frac{1}{s}$ (the true dimensions of charge) multiplied by $\frac{1}{s}$ (the omitted term), which produces $\frac{1}{s^2}$.

The same considerations apply to the size of the unit of this quantity. Since the charge is not defined independently of the equation, the fact that there is only one force involved means that the expression QQ' is actually $Q^{\frac{1}{2}}Q'^{\frac{1}{2}}$. It follows that, unless some structural factor (as previously defined) enters into the Coulomb relation, the value of the natural unit of Q derived from that relation should be the second power of the natural unit of $\frac{1}{s^2}$. In carrying out the calculation we find that a factor of 3 does enter into the equation. This probably has the same origin as the factors of the same size that apply to a number of the basic equations examined in Vol. I. It no doubt has a dimensional significance, although a full explanation is not yet available.

The natural unit of $\frac{1}{s^2}$, as determined in Vol. I, is 7.316889×10^{-6} sec/cm². On the basis of the findings outlined in the foregoing paragraphs, the value of the natural unit of charge is

$$Q = (3 \times 7.316889 \times 10^{-6})^2 = 4.81832 \times 10^{-10} \text{ esu.}$$

There is a small difference (a factor of 1.0032) between this value and that previously calculated from the Faraday constant. Like the similar deviation between the values for the gravitational constant, this difference in the values of the unit of charge is within the range of the secondary mass effects, and will probably be accounted for when a systematic study of the secondary mass relations is undertaken.

The equivalence of the scalar motions AB and BA, which plays an important part in the force relations, is also responsible for the existence of a unique feature of static electricity, the *induction* of charges. One of the basic characteristics of scalar motion, resulting from this equivalence is that it is indifferent to *location* in the reference system. From the vectorial standpoint, locations are very significant. A vectorial motion originating at location A and proceeding in the direction AB is specifically defined in the reference system, and is sharply distinguished from a similar motion originating at location B and proceeding in the direction BA. But since a scalar motion has magnitude only, a scalar motion of atom A toward atom B is simply a decrease in the distance between A and B. As such, it cannot be distinguished from a similar motion of B toward A. Both of these motions have the same magnitude, and neither has any other property.

Of course, the scalar motion *plus* the coupling to the reference system does have a specific location in that system: a specific reference point and a specific direction. But the coupling is independent of the motion. The factors that determine its nature are not necessarily constant, hence the motion AB does not necessarily continue on the AB basis. A change in the coupling may convert it to BA, or it may alternate between the two.

The rotational component of the scalar motion of a charged atom always maintains the same relation to an atom at another location. Half of the elements of that rotational motion are approaching the second atom, while the other half are receding in equivalent directions and at equivalent speeds. But this is not true of the rotational vibration that constitutes a charge. In this case the relation of the motion (charge) to the distant atom is continually changing; that is, the relative motion of the two atoms has the same vibratory character as the charge itself. As has been stated, a scalar motion A (such as a charge) toward or away from atom B is indistinguishable from a similar motion of B toward or away from A. The representation of this motion in the spatial reference system can therefore take either form.

Ordinarily, some redistribution of energy is required for a change from one representation of a motion to another, and such changes therefore do not usually take place in the absence of external forces. In fact, Newton's first law of motion requires a motion in the direction

AB to continue in that direction indefinitely unless acted upon by some force. However, there is an exception to this general rule because of the existence of a class of phenomena that we may call *zero energy processes*. Most of the physical processes that have been examined in the preceding pages either operate by application of energy, or occur spontaneously with release of energy. For instance, there is a force of cohesion between the atoms of a solid, and energy must be applied to separate them. If they are allowed to recombine, a corresponding amount of energy is liberated. But the various components of a combination of basic motions are not bound to each other in this manner in all cases. Often they are merely *associated*, and are free to separate or combine without gaining or losing energy. This is always recognized, even in macroscopic phenomena, in the case of translational motion. If this motion is transferred, as, for instance, from one billiard ball to another, it is understood that no expenditure of energy is required in order to detach the motion from the first ball. The full amount of energy separated from one object is received by the other, subject only to losses due to friction or other inefficiencies in the process itself.

Such transfers are zero energy processes. They are essentially accommodations to conditions in the environment. In the case of the billiard ball, the process is originated by the presence of the second ball in the line of travel. This makes a continuation of the motion in its original form impossible. All, or part, of the motion is therefore transferred to the ball that has an open path ahead of it, and *can* move. In the case of the charges, the zero energy process is originated by the existence of a force, that of the inducing charge, to which the components of a motion combination are able to respond.

The differences between these zero energy processes and those resulting from energy differentials play an important role in ionization of matter, and will be given further consideration in the next chapter.

One such zero energy process is the simultaneous creation or destruction of charges of the same magnitude and opposite polarity. It is the existence of this process, together with the equivalence of scalar motions AB and BA, that makes the induction of electric charges possible. As we saw earlier, all material objects contain a concentration of uncharged electrons, which are essentially rotating

units of space. In each case where an electron exists in an atom of matter, the atom likewise exists in the unit of space that constitutes the electron. This might be compared to a solution of alcohol in water. The atoms of alcohol exist in the water, but it is equally true that the atoms of water exist in the alcohol.

Let us now consider an example in which a positively* charged body X is brought into the vicinity of an otherwise isolated metallic object Y. The scalar direction of the vibratory motion (charge) of atom A in object X is periodically reversing, and at each reversal the reference point of the motion of A relative to any atom B that is free to move is redetermined by chance; that is, the motion may appear in the reference system either as a motion of A toward B or a motion of B toward A. By means of this chance process, the motion is eventually divided equally between AB and BA.

An atom C that is located in extension space is not free to move because energy would be required for the motion. But atom B in object Y, which is located in electron space, is not subject to this energy restriction, as the rotational motions of the atom and the associated electron are oppositely directed. The same motion that constitutes a positive* charge on the atom constitutes a negative* charge on the electron, because in this case it is related to a different reference point. The production of these oppositely directed charges is a zero energy process. It follows that atom B is free to respond to the periodic changes in the direction of the scalar motion originating at A. In other words, the positive* charge on atom A in object X induces a similar positive* charge on atom B in object Y and a negative* charge on the associated electron.

The electron is easily separable from the atom, and it is therefore pulled to the near side of object Y by the positive* charge on object X, leaving atom B in a unit of extension space, and with a positive* charge. The positions of the positively* charged atoms are fixed by the inter-atomic forces, and these atoms are not able to move under the influence of the repulsive forces exerted by charged object X, but the positive* charges are transferred to the far side of object Y by the induction process. The residual positive* charge on atom B induces a similar charge on a nearby atom D that is located in electron space. The electron at D, now with a negative* charge, is drawn to atom

B, where it neutralizes the positive* charge and restores the atom to the neutral status. This process is repeated, moving the positive* charge farther from object X in each step, until the far side of object Y is reached.

Where the original charge on object X is negative*, a negative* charge is induced on the electron associated with atom B. This is equivalent to a positive* charge on the atom. In this case, the negatively* charged electron is repelled by the negative* charge on object X and migrates to the far side of object Y. The residual positive* charge on the atom is then transferred to the near side of this object by the induction process.

If metallic object Y is replaced by a dielectric, the situation is changed, because in this case the electrons no longer have the capability of free movement. The induced charge on the atom and the opposite charge on the electron (or vice versa) remain joined. It is possible, however, for this atom to participate in a relative orientation of motions with a neutral atom-electron unit with which it is in contact, the result being a two-atom combination in which the negative* pole of one atom is neutralized by contact with the positive* pole of the other, leaving one atom-electron unit positively* charged and the other negatively* charged (that is, the charge is on the electron).

The optimum separation between the unlike charges, when under the influence of an external charge, the condition that is reached when the carriers of the negative* charges are free to move, is the maximum. The situation in the two-atom combination is therefore more favorable than that in the single atom, and the combination takes precedence. A still greater separation is achieved if one or more neutral atoms are interposed between the atoms of the charged combination. Each event of this kind moves either the positive* or the negative* charge in the direction determined by the inducing charge. Thus the effect of an inducing charge on a dielectric is a separation of the positive* and negative* charges similar to, but less complete than, that which takes place in a conductor, because the length of the chains of atoms is limited by thermal forces.

On the basis of the foregoing explanation, the charges are produced by induction. The subsequent separation is accomplished by action

of the inducing charge on the newly produced induced charges. Conventional theory of dielectrics is based on the concept of the nuclear atom, a hypothetical structure in which the components are held together by the attraction between positive* and negative* charges. It is assumed that these charges have a limited amount of freedom of movement, and can separate slightly on being subjected to the effect of an external charge. One observation that has been interpreted as supporting the assumption that pairs of positive* and negative* charges are always present in the atoms is that if a charged dielectric is subdivided, each of the parts contains both positive* and negative* charges. This is quite different from the behavior of charges in conductors. If a metallic object is cut perpendicular to the line of force while under the influence of an inducing charge, the two parts are oppositely charged, and will remain so after the inducing charge is removed. But if the same procedure is followed with a dielectric, both parts have positive* and negative* charges on the opposite sides, just as in the original object before separation. And when the inducing charge is removed, both parts revert to the neutral status. The current interpretation of these results, as expressed in a contemporary textbook, is this:

The inference to be drawn is that insulators contain charges which can move small distances so that attraction still occurs, but that they are bound in equal and opposite amounts so that no splitting of the body can separate two kinds of charge.⁵⁷

The amount of separation of charges that could take place in the manner assumed by this theory is admittedly very small, and it is difficult to account for the generation of any substantial attractive or repulsive forces by this means. But forces of this nature actually do exist. Small static charges, usually produced by friction, are common in the terrestrial environment, and they produce effects that are quite noticeable. Merely walking across a carpeted room in cold, dry weather can build up enough charge to give one an uncomfortable sensation when he touches a metallic object and discharge occurs. Likewise, the behavior of the modern synthetic fabrics shows the effect of static charges, including the inductive effect, in a conspicuous, and often annoying, way. These fabrics

behave in much the same way as charged conductors. They attract such things as bits of paper and chips of wood, and are themselves attracted by the furniture and walls of a room.

The discrepancy between the very small theoretical separation of charges and the relatively large inductive effect has forced the theorists to call upon collateral factors, such as the presence of contaminants, to explain the observations. For example, the following statement taken from a physics textbook refers to the ability of electrically charged non-conducting objects to pick up bits of paper and wood:

A chip of perfect insulator would show hardly any effect, but bits of wood and paper always have enough moisture to make them slightly conducting.⁵⁸

Relying on a collateral process of this kind when a straightforward explanation cannot be found is always a kind of “last resort” expedient. In this case the “moisture” hypothesis also conflicts with the well-known tendency of dry conditions to *increase* the static effects. It is therefore significant that the explanation that the Reciprocal System of theory provides for the ability of static charges to attract or repel dielectrics is based on the electrical properties of the dielectric materials, not on contaminants. The much greater separation of charges that results from the inductive process described in this chapter resolves this problem, while it remains consistent with the appearance of charges at both ends of each piece when a dielectric under the influence of an inducing force is separated. Before the separation takes place a substantial number of atoms of the dielectric exist in multi-atom combinations with positively* and negatively* charged ends. Although the separation of the charges in many of these combinations is large compared to the distance between atoms, it is very small compared to the dimensions of an ordinary charged dielectric. Thus when the separation occurs, there are charged combinations of this kind in each portion. Consequently, each piece has the same charge characteristics as the original unbroken object.

It was pointed out in Volume I that the existence of positive* and negative* charges in close proximity, as required by the nuclear theory

of the atom, is incompatible with the observed behavior of charges of opposite polarity. These observations show that such charges neutralize each other long before they reach separations as small as those which would exist in the hypothetical nuclear atom. This is a decisive argument against the validity of the nuclear theory. It is appropriate, therefore, to note that the existence of both positive* and negative* charges in objects under the influence of inducing charges does not conflict with our finding that there is a minimum distance (identified as the natural unit of distance, 4.56×10^{-6} cm) within which charges of opposite polarity cannot coexist. The coexistence of induced positive and negative charges is possible because they are forcibly prevented from reaching the limiting distance at which they would combine. If the external charge is removed, the induced charges do combine and neutralize each other.

In charging by induction it is often convenient to make use of *grounding*, which is simply connecting the inductively charged object to the earth by means of a conductor. The earth is electrically neutral, and so large that it is insensitive to gains or losses of charge in the amounts actually encountered in practice. If object Y is grounded while under the influence of a negative* inducing charge, the negatively* charged electrons on the far side of the object are forced through the conductor into the earth. Breaking the ground connection then leaves only positive* charges on object Y, and this object remains positively* charged after the object X that contains the inducing charge is removed. If the induction process is initiated by a positive* charge on object X, the ground connection permits electrons to be pulled from the earth and charged negatively* to neutralize the positive* charges on Y, leaving only negative* charges. Breaking the ground connection then leaves Y negatively* charged.

The locations occupied by the charges in any charged conducting object not subject to inductive forces are determined by the repulsion between the charges, which operates to produce the maximum separation. If the object is under the influence of outside charges, the charge locations are determined by the net effect of the inductive potential and the repulsion between like charges. In either case, the result is that the charges are confined to the outside surfaces of the conducting materials, and, except for local variations in very irregular

bodies, there are no charges in the interiors. The same considerations apply if the objects are hollow. The inside walls of such objects carry no charges. These walls may be charged by placing an insulated charged object in the hollow interior, but in that case the inside walls are “outside” from the standpoint of the inducing charge; that is, they are the locations closest to that charge.

The observed concentration of charge at the conductor surfaces is another direct contradiction of the accepted theory of the electric current, which views the current as a movement of charges. This concentration at the surface is due to the mutual repulsion between the charges, which drives them to opposite sides of the conductor. The repulsive force is not altered if the charges move along the conductor, since the direction of this force is perpendicular to the direction of movement. Nor would the presence of positive* charges on the interior atoms of the conductor alter this situation, if any such charges existed. If the electrons, or any portion of them, are firmly held by the attraction of the hypothetical proton charges, then they cannot move as an electric current. If they are free to move in response to an electric potential gradient, then they are also free to move to the surfaces of the conductor under the influence of their mutual repulsions.

From this it follows that if present-day electric theory were correct the current should flow only along the outer surfaces of the conductors. In fact, however, electric resistance is generally proportional to the cross-sectional area of the conductor, indicating that the motion takes place fairly uniformly throughout the entire cross section. This adds one more item of evidence supporting the finding that the electric current is a movement of uncharged electrons, not of charges.

Since no charges are induced within a hollow conductor by an outside charge, any object within a conducting enclosure is insulated against the effects of an electric charge. Similar elimination, or reduction, of these effects is accomplished by conductors of other shapes that are interposed between the charge and the objects under consideration. This is the process known as *shielding*, which has a wide variety of applications in electrical practice.

Within the limits to which the present examination of electrical

phenomena has been carried, there does not appear to be any major error in the conventional dimensional assignments, other than those discussed in the preceding pages. Aside from the errors that have been identified, the SI system is dimensionally consistent, and consistent with the mechanical system of quantities. The space-time dimensions of the most commonly used electrical units are listed in Table 28. The first column of this table lists the symbols that are used in this work. The other columns are self-explanatory.

TABLE 28: *ELECTRIC QUANTITIES*

t	time	second	t
	dipole moment	coulomb(t/s) $\times$ meter	t
W	energy (work)	watt-hour	t/s
Q	charge (flux)	coulomb (t/s)	t/s
V	potential	volt	t/s ²
V	voltage	volt	t/s ²
E	field intensity	volt/meter	t/s ³
	flux density	coulomb (t/s)/meter ²	t/s ³
	charge density	coulomb (t/s)/meter ³	t/s ⁴
	resistivity	ohm-meter	t ² /s ²
R	resistance	ohm	t ² /s ³
	current density	ampere/meter ²	1/st
	power	watt	1/s
D	displacement	coulomb (s)/meter ²	1/s
P	polarization	coulomb (s)/meter ²	1/s
S	space	meter	s
Q	electric quantity	coulomb	s
C	capacitance	farad	s
I	current	ampere	s/t
ϵ	permittivity	farads/meter	s ² /t
σ	conductivity	siemens/meter	s ² /t ²
	conductance	siemens	s ³ /t ²

The natural units of most of these quantities can be derived from the natural units previously evaluated. Those of the remaining quantities can be calculated by the methods used in the previous determinations,

but the evaluation is complicated by the fact that the measurement systems in current use are not internally consistent, and it is not possible to identify a constant numerical value that relates any one of these systems to the natural system of units, as was done for the mechanical quantities that involve the unit of mass. Neither the SI system nor the cgs system of electrical units qualifies as a single measurement system in this sense. Both are combinations of systems. In the present discussion we will distinguish the measurement systems, as here defined, by the numerical coefficients that apply, in these systems, to the natural unit of space, s , and inverse speed, $\frac{1}{s}$.

On the basis of the values of the natural units of space and time in cgs terms established in Volume I, the numerical coefficient of the natural unit of s , regardless of what name is applied to the quantity, should be 4.558816×10^{-6} , while that of the natural unit of $\frac{1}{s}$ should be 3.335635×10^{-11} . In the mechanical system of measurement the quantity s is identified in its most general character as space, and the unit has the appropriate numerical coefficient. But the concept of mass was introduced into the $\frac{1}{s}$ quantity, here called energy, and an arbitrary mass unit was defined. This had the effect of modifying the numerical values of the natural units of energy and its derivatives by the factor 4.472162×10^7 , as explained in Volume I.

The definition of the unit of charge (esu) by means of the Coulomb equation in the electrostatic system of measurement was originally intended as a means of fitting the electrical quantities into the mechanical measurement system. But, as pointed out in Chapter 14, there is an error in the dimensional assignments in that equation which introduces a deviation from the mechanical values. The electrostatic unit of charge and the other electric units that incorporate the esu therefore constitute a separate system of measurement, in which $\frac{1}{s}$ is identified with electric charge. The unit of this quantity was evaluated from the Faraday constant in Chapter 9 as 4.80287×10^{-10} esu.

Unit charge can also be measured directly, inasmuch as some physical entities are incapable of taking more than one unit of electric charge. The charge of the electron, for instance, is one unit. Direct measurement of this charge is somewhat more difficult than the derivation of the natural unit from the Faraday constant, but the direct measurements

are in reasonably good agreement with the indirectly derived value. In fact, as noted in Chapter 14, clarification of the small scale factors that affect these phenomena will probably bring all values, including the one that we have derived theoretically, into agreement.

An electromagnetic unit (emu) analogous to the esu can be obtained by magnetic measurements, and this forms the basis of an electromagnetic system of measurement. The justification for using the emu as a unit of electrical measurement is provided by the assumption that it is an electric unit derived from an electromagnetic process. We now find, however, this it is, in fact, a magnetic unit; that is, it is a two-dimensional unit. It is therefore a unit of t^2/s^2 rather than a unit of t/s . To obtain an electric (one-dimensional) unit, t/s , corresponding to the esu from the emu it is necessary to multiply the measured value of the emu coefficient, 1.602062×10^{-20} by the natural unit of s/t , 2.99793×10^{10} cm/sec. This brings us back to the electrostatic unit, 4.80287×10^{-10} esu. The electromagnetic system is thus nothing more than the electrostatic system to which an additional factor, meaningless in the electrical context, has been applied.

The SI system of units is a modification of the electromagnetic system. In the early days of electrical measurement the ampere was selected as the fundamental unit, and was defined on an arbitrary basis. After more information had been accumulated, and the desirability of relating the measurement system to physical fundamentals was recognized, the electromagnetic (emu) system was adopted for general use, but in order to avoid making a radical change in the size of the ampere, an arbitrary factor of 10 was introduced. As M. McCaig remarks, the appearance of such a number “in a primary definition is unusual; it arises because although this definition is intended to fix the value of the ampere, we have already decided in advance fairly precisely the value we desire the unit to have.”⁵⁹

The arbitrary modification of the emu values changed the numerical coefficient of the natural unit of t/s to 1.602062×10^{-19} . Because of the lack of distinction between electric charge (t/s) and electric quantity (s) in current practice, the same unit is used for both of these physical quantities in all three of the electrical measurement systems, as shown in Table 29.

TABLE 29: *NUMERICAL COEFFICIENTS OF THE NATURAL UNITS*

	s	$\frac{t}{s}$
Space-time (cgs)	4.558816×10^{-6}	3.335635×10^{-11}
Mechanical	4.558816×10^{-6}	1.49175×10^{-3}
Electrostatic	4.80287×10^{-10}	4.80287×10^{-10}
Electromagnetic	1.602062×10^{-20}	1.602062×10^{-20}
SI modification	1.602062×10^{-19}	1.602062×10^{-19}

In applying the principle of the equivalence of natural units to electrical quantities it is necessary to take into account these differences between the numerical values applying to the different systems. For example, the natural unit of capacitance, the quantity that plays the principal role in the phenomena discussed in Chapter 15, is the natural unit of electric charge divided by the natural unit of voltage, $\frac{t}{s} \times s^2/t = s$. On the basis of the explanation of the natural electrical units given in the preceding paragraphs, the value of the natural unit of electric charge in the cgs electrostatic system is 4.80287×10^{-10} esu. The natural unit of capacitance is this value divided by the natural unit of voltage, which was evaluated in Chapter 9 as 9.31146×10^8 volts. The result is 5.15802×10^{-18} farads. As we found earlier, the farad is a unit of space. The natural unit of space derived in Volume I is 4.558816×10^{-6} cm. Dividing these two values, we obtain 1.1314×10^{-12} as the ratio of the numerical coefficients of the natural units. From geometric measurements, the centimeter, as a unit of capacitance, has been found equal to 1.11126×10^{-12} farads. The theoretical and experimental values are therefore in agreement within the limits of accuracy to which the present study of the electric relations has been carried.

In this case the application of the equivalence principle merely corroborates an experimental result. Its value as an investigative tool derives from the fact that it is equally applicable in situations where nothing is available from other sources.

CHAPTER 17

Ionization

Electric charges are not confined to electrons. Units of the rotational vibration that constitutes electric charge may also be imparted to any other rotational combination, including atoms as well as other sub-atomic particles. The process of producing such charges is known as *ionization*, and electrically charged atoms or molecules are called *ions*. Like the electrons, atoms or molecules can be charged, or ionized, by any of a number of agencies, including radiation, thermal motion, other physical contact, etc. Essentially, the ionization process is simply a transfer of energy, and any kind of energy will serve the purpose if it is delivered to the right place and in the necessary concentration.

As indicated above, one of the sources from which the ionization energy can be derived is the thermal energy of the ionizable matter itself. We saw in Chapter 5 that the thermal motion is always directed outward. It therefore joins with ionization in opposition to the basic inward rotational motions of the atoms, and is to some degree interchangeable with ionization. The magnitude of the energy required to ionize matter varies with the structure of the atom and with the existing level of ionization. Each element therefore has a series of ionization levels corresponding to successive units of rotational vibration. When the thermal energy concentration (the temperature) of an aggregate reaches such a level the impacts to which the atoms are subjected are sufficiently energetic to cause some of the linear thermal motion to be transformed into rotational vibration, thus ionizing some of the atoms. Further rise in temperature results in ionization of additional atoms of the aggregate, and in additional ionization (more charges on the same atoms) of previously ionized matter.

Thermal ionization is only of minor importance in the terrestrial environment, but at the high temperatures prevailing in the sun and

other stars thermally ionized atoms, including positively* charged atoms of Division IV elements, are plentiful. The ionized condition is, in fact, normal at these temperatures, and at each of the stellar locations there is a general *ionization level* determined by the temperature. At the surface of the earth the electric ionization level is zero, and except for some special cases among the sub-atomic particles, any atom or particle that acquires a charge while in the gaseous state is in an unstable condition. It therefore eliminates the charge at the first opportunity. In some other region where the prevailing temperature corresponds to an ionization level of two units, for example, the doubly ionized state is the most stable condition, and any atoms that are above or below this degree of ionization tend to eliminate or acquire charges to the extent necessary to reach this stable level.

Since the rotational vibration that we know as ionization is basically a motion in opposition to the rotational motion of the atom, the ionization cannot exceed the net effective positive* displacement (the atomic number). In a region where the ionization level is very high, the heavier elements therefore have a considerably larger content of positive* displacement in the form of ionization at a given temperature than those of smaller mass. This point has an important bearing on the life cycle of the elements, and will be given further consideration later.

In the nuclear theory of atomic structure currently accepted by the physicists the atomic “nucleus” is surrounded by a number of electrons equal to the atomic number of the element. Ionization is viewed as a process of detaching electrons from the atom. On this basis, the maximum degree of ionization is attained when all electrons have been removed and only the bare nucleus remains. This is a plausible hypothesis, and, on first consideration, its plausibility would appear to be a point in favor of the nuclear theory. It should be realized, however, that any tenable theory of atomic structure will have essentially the same explanation of ionization, differing only in the language in which it is expressed. Such a theory must identify entities that are added to, or removed from, the atom as the atomic number increases. Successive addition or elimination of these entities then explains ionization. In the nuclear theory, which views the atom as a collection of particles, these entities are electrons. In

the theory of the universe of motion, which finds the atom to be a combination of motions, they are units of rotational motion. Any other theory that might be formulated would necessarily have to identify some entity that could similarly be added or removed unit by unit. Thus the ionization process would be consistent with any theory. Consequently, it gives support to none.

In the terrestrial environment each ionization level of each element has a specific *ionization potential* that represents the amount of energy required in order to accomplish the ionization. It is currently assumed that these values are fixed natural relations and therefore constant for all environments. The theoretical status of this assumption in the context of the Reciprocal System of theory has not yet been clarified. It may well be valid throughout the gaseous state. However, the measured ionization levels are obviously not applicable to ionization in the condensed gas state, the state in which the gas molecules are within the equivalent of unit distance of each other. The physical relations in this state are very different from those in an ordinary gas, including reversal of all scalar directions. Thus all that we can now say about the ionization potential in this state is that each successive level of ionization must involve an increase in energy. As we will see in Volume III, the matter in most of the observed stars is in the condensed gas state.

The relation between temperature and the degree of ionization enables using the ionization, which can be observed spectroscopically, as a measure of the surface temperature of the stars. For example, below 12,000 K, helium is not ionized. At about 35,000 K it is mainly in the form of He II (singly ionized). At still higher temperatures it is doubly ionized (He III). Other elements have similar ionization patterns, and the mixture of ions observed in the spectrum of a star thus indicates the range of temperature at its surface. The O stars, which are in the range up to about 80,000 K are reported to contain N II, O II, C II, and Si III, as well as helium and hydrogen ions.

It should be understood, however, that this relation between ionization and temperature holds good only where *the ionization is produced thermally*. References are made in the astronomical literature to “ionization temperatures,” but these are merely the temperature equivalents of the ionization levels. Unless the ionization

is thermally produced they do not indicate the actual temperature. The level of ionization is a reflection of the strength of the ionizing agency, whatever it may be. If that agency is the thermal energy, then the ionization is a measure of the temperature. But if the ionizing agency is radiation, the ionization level is a measure of the strength of the radiation, not the temperature.

In Volume III we will encounter the same kind of a misconception in dealing with the relation between temperature and the production of x-rays. When the x-rays are thermally produced, there is actually a relation between the x-ray emission and the temperature, but here, again, if the x-rays are produced by some other agency, the relation is between the x-ray emission and the strength of that other agency, and it is independent of the temperature. The importance of this point lies in the fact that the emission of x-rays is currently being treated as an indication of high temperature in cases where the nature of the x-ray production process is unknown; even in cases where the conditions are such that the temperatures necessary for thermal production of x-rays are impossible. Temperatures in the millions of degrees are inferred from x-ray observations in locations where the actual temperature level cannot be more than a few degrees above absolute zero.

“Temperature,” without a qualifying adjective, is a specifically defined concept, and it is temperature *as thus defined* that enters into the various thermal relations. The use of other kinds of “temperature” is entirely in order, providing that each is clearly defined, and is identified by an appropriate adjective, in an expression such as “*ionization temperature*.” In fact, we will introduce such an alternate kind of temperature, a *magnetic* temperature, in Chapter 24. But it should be recognized that these “temperatures” have their own sets of properties. The thermal relations do not apply to them. For example, the general gas law applies only to temperature in the usual (thermal) sense. This law is expressed as $PV = RT$, where P is the pressure, V is the volume, T is the temperature, and R is the gas constant. From this law it is apparent that a high temperature can be developed in a given volume of gas only under high pressure. In interstellar and intergalactic space the pressure acting on the extremely tenuous medium is near zero, and from the general gas law

it is evident that the temperature must be at a correspondingly low level. The temperatures in the millions of degrees that are currently being reported from these regions are totally unrealistic, if they are intended to mean "temperature" in the thermal sense.

Some of the existing confusion in this area appears to be due to a failure to draw a clear distinction between the two types of vectorial motion in which the particles of a gas participate. These constituent particles share in the translational motion of a gaseous aggregate as a whole, and it is generally understood that this is not a thermal motion; that is, a fast-moving aggregate may be relatively cool. An atom or particle moving independently in space is subject to the same considerations. Its free translational motion has no thermal significance. The thermal motion is a product of containment. It is the directionally distributed random motion that results from the restriction of the motion to the volume within certain limits. The pressure is a measure of the *containment*. The temperature, the measure of the thermal motion, is therefore a function of the pressure, as indicated in the gas laws. High temperatures can only be attained under high pressures. If part, or all, of the gas in an aggregate escapes from confinement, its constituents move outward unidirectionally, and the thermal motion is converted to linear translational motion. The temperatures and pressures decrease accordingly.

The picture of the nature of electric charges and ionization that we derive from the postulates of the theory of the universe of motion is very different from the currently accepted explanation of these phenomena, which is an outgrowth of hypotheses formulated in the early days of electrical investigation on the basis of the limited amount of empirical information then available. The early investigators in this area identified negative* charges with electrons and positive* charges with atoms of matter. Meanwhile it was found that the atoms of certain elements undergo spontaneous disintegration in which electrons are emitted along with other products. On the basis of these empirical findings, the scientific community adopted the hypothesis previously mentioned in which positive* charges are attributed to an atomic "nucleus," and negative* charges entirely to electrons. Positive* and negative* ionizations were then ascribed to deficiency or excess of electrons, respectively.

One disturbing feature of this explanation was the great disparity in the sizes of the units of the two entities that were identified as the carriers of the charges. The roles to be played by positive* and negative* charges in the theory were essentially reciprocal in nature, yet the presumed carrier of the positive* charge, the proton, has nearly two thousand times the mass of the negatively* charged particle, the electron. Physicists were therefore greatly relieved when the positive* analog of the electron, the positron, was discovered. It does not seem to be generally appreciated that this discovery, which restored the symmetry that we have come to expect in nature, has destroyed the foundations of the orthodox theory. It is now evident that the positive* charge is as much of a reality as the negative* charge; it is not merely an electron deficiency, as the theory contends.

While the discovery of the positron solved one of the symmetry problems, it produced another that has been even more troublesome. Inasmuch as the electron and the positron are inversely related, so far as we can tell, it would seem that they should appear in equal numbers. But positrons are scarce in our environment, whereas electrons are plentiful. Conventional science has no answer to this problem, other than mere speculations. From the theory of the universe of motion we find that the asymmetrical distribution of electrons and positrons, and of positive* and negative* charges in general, is not due to any inherent difference in the character of the motions that constitute the charges, but is a consequence of the fact that the net rotational displacement of the atoms of ordinary matter is in time; that is, it is positive. The charges acquired by these atoms in the ionization process are therefore positive*, except in the relatively few instances where negative* ionization is possible because of the existence of negative electric rotational displacement of the appropriate magnitude in the structures of certain atoms. The simple positively* charged sub-atomic particles, the positrons, are scarce in the vicinity of material atoms because their net rotational time displacement is compatible with the basic structure of the atoms, and they are readily absorbed on contact. The corresponding negatively* charged particles of the material system, the electrons, are abundant, as their space displacement is usable in the structures of the material atoms only to a very limited degree.

It is evident that both of the mechanisms discussed in the foregoing pages, the selective incorporation of the positrons into the structure of matter, which leaves a surplus of free electrons, and the ionization mechanism, which produces only positive* ions under high temperature conditions (where most of the ionization takes place), are incompatible with the existence of a law requiring absolute conservation of charge. This will no doubt disturb many individuals, because the conservation laws are generally regarded as firmly established basic physical principles. Some consideration of this issue will therefore be appropriate before moving on to other subject matter.

In conventional physical science the conservation laws are empirical. As expressed by one physicist:

We are in a curious situation. We know the conservation laws, but we do not know their underlying dynamic basis; that is, we do not know the kind of symmetries responsible for them.⁶⁰

While the conservation laws have retained their original status as important fundamental principles of physics during the broad expansion of scientific knowledge that has taken place in the twentieth century, the general understanding of their nature has undergone a significant change. Any empirically based relation or conclusion is always subject to modification by reason of relevant new discoveries. This is what has happened to conservation. Originally, the law of conservation of energy, for instance, was thought to be inviolable. "No gain or loss of energy has ever been observed in an isolated system," says a 1919 textbook.⁶¹ This statement is no longer true. Mass and energy, we have found, are interconvertible. Thus one can increase at the expense of the other. The content of a conservation law has therefore had to be redefined. As expressed by Eric M. Rogers,

In its present fullest form you may consider it [the conservation of energy] more than a generalization from experiment; it has expanded into a convention, an agreed scheme of energy now so defined that its total must, by definition, remain constant.⁶²

It is now frequently stated that we should not speak of the conservation of mass or the conservation of energy, only the conservation of mass-energy. However, the conversion of one of these entities into the other

occurs only under circumstances that, in the terrestrial environment, are quite exceptional, and the separate conservation laws are applicable under all ordinary circumstances. It would seem more appropriate, therefore, to state these laws individually, as in the past, and to qualify the statements in such a way as to limit the application of the laws to situations in which there is no conversion to or from a different form of motion.

These same considerations apply to electric charges. There is a wide range of physical activity in which the conservation of charge is maintained. Indeed, the currently prevailing view is that charge conservation is absolute, as indicated in the following statement:

The law of conservation of electric charge...states that...there is no way to alter, to the slightest degree, the total amount of electric charge in the world.⁶³

Our finding is that all physical quantities with the dimensions $\frac{1}{s}$, including electric charge, are equivalent to, and, under appropriate conditions, interconvertible with kinetic energy. Thus while energy and charge are each conserved individually within a certain range of physical processes, there is a wider range of processes in which the quantity $\frac{1}{s}$ is conserved, but changes occur in the magnitudes of the subsidiary quantities, such as charge or kinetic energy, because of conversion from one to another.

The law of conservation of electric charge is valid wherever no such conversion takes place, and it has persisted because most of the common electrical processes are of this nature. The observation that has been most influential in leading to the conclusion that charge conservation is absolute is the existence of processes in which positive* and negative* charges are created in pairs, and destroyed jointly. A unit negative* charge is a unit of outward scalar motion in time. A unit positive* charge is a unit of outward scalar motion in space. Since the two motions are oppositely directed from the natural zero point, a combination of the two units arrives at a net total motion (measured as energy or speed) of zero on the natural scale. Thus the creation or neutralization of such a pair of charges involves no change in the total net charge or energy. It is another instance of what we have called a zero energy process.

The induction process discussed in Chapter 16 is another example. As explained there, an external positive* charge induces a rotational vibration (charge) which is positive* relative to each of the atoms of the object subjected to the charge, and negative* relative to the mobile units of space (electrons) in which some of these atoms are located. The attractive and repulsive forces due to the external charge then cause each of the atom-electron combinations to separate into a pair of positively* and negatively* charged entities. It can be seen that this process does not alter the net amount of electric charge. An object (a combination of motions) with zero net rotational vibration (charge) separates into two components, the net total charge of which is zero.

However, it is also evident that these are processes of a special kind, and the fact that charge is conserved in such processes does not indicate that charge is always conserved. The best resolution of the conservation question appears to be to recognize that each of the conservation laws previously formulated is valid within certain limits, and therefore has a specific field of usefulness, but to state each of these laws in such a form that its applicability is restricted to the range of conditions in which no conversion from or to other forms of motion is involved.

While the foregoing is a significant limitation of the field of applicability of the charge conservation law, there is still a wide range of physical phenomena in which electric charge is conserved, as the processes that involve changes in the net total $\frac{1}{2}$ in the form of electric charge are confined mainly to those that take place at very high temperatures, or very large kinetic energies.

One of the important areas in which charge is conserved is ionization in liquids. The molecules of a simple chemical compound such as hydrochloric acid (HCl), for example, consist of two components, in this case a hydrogen atom and a chlorine atom, oriented in the manner described in Volume I, and held together by the cohesive forces discussed in Chapter 1 of this volume. In the liquid state the molecules move independently, subject to the restrictions imposed by the nature of this state of matter. The effective rotation of the hydrogen atom, as oriented in HCl, is positive, while that of the chlorine atom is negative. These components of the molecule are therefore capable of taking positive* and negative* charges respectively, if they separate.

The molecules in a liquid aggregate are in constant motion, and collisions are frequent. A certain percentage of these collisions, depending on the temperature, are energetic enough to break the bonds between the molecular components and separate each molecule into two parts. Ordinarily these parts recombine promptly, but if the atom is located in a unit of electron space, the collision imparts a rotational vibration to each of the components. (As noted in Chapter 13, such rotational vibrations, electric charges, are easily produced in contacts of various kinds.) This rotational vibration is a positive* motion of the hydrogen atom relative to the associated electron space, and a negative* motion of the electron relative to the chlorine atom. The generation of the charges is thus a zero energy process, and it does not add to the energy of the system.

The HCl molecule has now become a H^+ molecule, an *ion*, and a Cl atom associated with a charged electron, a Cl^- ion, we may say. The charges on these new molecules, or ions, balance the valences of their associated atoms, and the ions are therefore stable in the same sense as the original HCl molecules, except that there is a rather strong tendency toward recombination that limits the net amount of ionization.

Let us now turn to an examination of the effects that are produced when a voltage is applied in such a way as to cause a voltage gradient in a liquid that is, to some extent, ionized. This is accomplished by inserting two electrical conductors, or *electrodes*, into the liquid, and connecting them through a source of current so that electrons are withdrawn from the positive* electrode, the *anode*, and forced into the negative* electrode, the cathode. Liquids such as HCl are not conductors of electricity, in the sense in which this term is applied to metals; that is, they do not permit free movement of electrons. However, the introduction of a voltage differential causes a movement of the *ions* in the ionized liquid.

As we saw in Chapter 15, this voltage differential forces some of the electrons at the cathode out into the spatial equivalent of time, and withdraws a similar number of electrons from the spatial equivalent of time at the anode. Some of the contacts with liquid molecules are sufficiently energetic to impart charges to electrons in the vicinity of the cathode. Thus a quantity of negative* charge accumulates in the liquid in this vicinity, a process known as *polarization*.

At the anode, the withdrawal of electrons leaves a deficiency of electrons, relative to the equilibrium concentration. This leads to a break-up of some of the neutral combinations of positive* atoms and negative* electrons. The electrons thus released are absorbed into the electron "vacuum," losing their charges in the process. This leaves a surplus of positively* charged ions; that is, the region in the vicinity of the anode is positively* polarized.

As a result of the polarization, the positive* and negative* ions are attracted to the cathode and anode respectively by the electric forces between unlike charges. The positive* ions (such as H^+) arriving at the cathode neutralize negatively* charged electrons, and withdraw them from the electron concentration in equivalent space. These are replaced by electrons drawn from the cathode. Additional electrons then acquire charges by the collision process to restore the polarization equilibrium in the liquid surrounding the cathode. Meanwhile the negative* ions (such as Cl^-) arriving at the anode neutralize positive* charges in the vicinity of that electrode, and release electrons, which are drawn into the anode to restore the polarization equilibrium.

The loss of electrons from the cathode and acquisition of electrons by the anode in the process that has been described creates a voltage difference between the two electrodes, in addition to that supplied by the external voltage source. A current therefore flows from the anode to the cathode through the metallic conductor to restore the equilibrium condition. This current persists as long as the ions continue to move through the liquid.

The proportion of the total number of molecules that will be ionized in a particular liquid under specified conditions is a probability function, the value of which depends on a number of factors, including the strength of the chemical bond, the nature of the other substances present in the liquid, the temperature, etc. Where the bond is strong, as in the organic compounds, the molecules often do not ionize at all within the range of temperature in which the substance is liquid. Substances such as the metals, in which the atoms are joined by positive bonds, likewise cannot be ionized in the liquid state, since the zero energy ionization process depends on the existence of a positive*-negative* combination.

The presence or absence of ions in the liquid is an important factor in many physical and chemical phenomena, and for that reason chemical compounds are often classified on the basis of their behavior in this respect as polar or non-polar, electrolytes or non-electrolytes, etc. This distinction is not as fundamental as it might appear, as the difference in behavior is merely a reflection of the relative bond strength: whether it is greater or less than the amount necessary to prevent ionization. The position of organic compounds in general as non-electrolytes is primarily due to the extra strength of the two-dimensional bonds characteristic of these compounds. It is worth noting in this connection that organic compounds such as the acids, which have one atom or group less strongly attached than is normal in the organic division, are frequently subject to an appreciable degree of ionization.

Ionization of a liquid is not a process that continues to completion; it is a dynamic equilibrium similar to that which exists between liquid and vapor. The electric force of attraction between unlike ions is always present, and if an ion encounters one of the opposite polarity at a time when its thermal energy is below the ionizing level, recombination will occur. This elimination of ions is offset by the ionization of additional molecules whose energy reaches the ionizing level. If conditions are stable, an equilibrium is reached at a point where the rate of formation of new ions is equal to the rate of recombination.

The conventional explanation of the ionization process is that it consists of a transfer of electrons from one atom, or group of atoms, to another, thus causing a deficiency of electrons, identified as a positive* charge, in one of the participants, and an excess of electrons, identified as a negative* charge, in the other. In the electrolytic process, the negative ions are assumed to carry electrons to the anode, where they leave the ions, enter the conductor, and flow through the external circuit to the cathode. Here they encounter the positive* ions that have been drawn to this electrode, and the charges are neutralized, restoring the electrical balance.

This is a simple and plausible explanation. It is not surprising, therefore, that it has met with widespread acceptance. Like many

another attractive, but erroneous, hypothesis, however, its net effect has been to direct physical thinking into unproductive channels. In fact, this interpretation of the electrolytic process is one of the major influences contributing to the belief that the electric current is a movement of charges, one of the basic errors of present-day electrical theory.

Since negative* charges clearly do move through the electrolyte to the anode, there is, on first consideration, an analogy with the metallic circuit, and discussions of electrolysis habitually refer to "passing a direct current through an electrolytic solution." If there actually were a continuous flow around the circuit, and if the moving units could be identified as negative* charges in one segment of that circuit, it would be reasonable to assume that the moving units in the remainder of the circuit are also charges. But this argument is wholly dependent on the continuity, and that continuity clearly does not exist. The electrolytic process is not a simple flow of current around the circuit; it is a more complex series of events in which both positive* and negative* charges originate in the solution and move in opposite directions to the electrodes. This means that electrolytic conduction has to be explained independently of metallic conduction, and it eliminates most of the support that the electrolytic process has been presumed to give to the conventional theory of the electric current.

The final topic for consideration in this chapter is the overall limit on the magnitude of the combined thermal and ionization energy. As pointed out earlier, the thermal energy must reach a certain level, which depends on the characteristics of the atoms involved, before thermal ionization is possible. After this level is reached, an equilibrium is established between the temperature and the degree of ionization. A further increase in the temperature of an aggregate causes both the linear speed displacement (particle speed) and the charge displacement (ionization) to increase, up to the point at which all of the elements in the aggregate are fully ionized; that is, they have the maximum number of positive* charges that they are capable of acquiring. Beyond this point of maximum ionization a further increase in the temperature affects only the particle speeds. Ultimately the total of the outward displacements (ionization and thermal) reaches equality with one of the inward magnetic rotational displacement

units of the atom. The inverse speed displacements then cancel each other, and the rotational motions that are involved revert to the linear status. At this point the material aggregate has reached what we may call a *destructive limit*.

There have been many instances in the preceding pages in which a limiting magnitude of the particular physical quantity under consideration has been shown to exist. We have just seen that the number of units of electric ionization of an atom is limited to the net equivalent number of units of effective electric rotational displacement. For example, the element magnesium, which has the equivalent of 12 net effective electric rotational displacement units, can take 12 units of electric vibrational displacement (ionization), but no more. Similarly, we found that the maximum rotational base of the thermal vibration in the solid state is the primary magnetic rotation of the atom. Most of the limits thus far encountered have been of this type, which we may designate as *non-destructive limits*. When such a limit is reached, further increase of this particular quantity is prevented, but there is no other effect.

We are now dealing with a quantity, the total outward speed displacement, which is subject to a different kind of a limit, a destructive limit. The essential difference between the two stems from the fact that the non-destructive limits merely define the extent to which certain kinds of additions to, or modifications of, the constituent motions of the atoms can be carried. Reaching the electric ionization limit only means that no more units of positive* electric charge can be added to the atom; it does not, in any way, imperil the existence of the atom. On the other hand, a limit that represents the attainment of equality with a basic motion of the atom has a deeper significance. Here it should be remembered that rotation is not a property of the scalar motion itself; it is a property of the coupling of the motion to the reference system. For example, the basic constituent of the uncharged electron is a unit of inward scalar motion in space. This motion *per se* has no properties other than the unit inward magnitude, but it is coupled to the reference system in such a way that it becomes a rotation, in the context of that system, retaining its inward scalar direction. When the electron is charged, the coupling is so modified that an oppositely directed

rotational vibration is superimposed on the rotation. The charged positron is a unit inward motion in time, similarly coupled to the reference system.

When brought into proximity, a charged electron and a charged positron are attracted toward each other by the electrical forces. When they make contact, the two rotational vibrations of equal magnitude and opposite polarity cancel each other. The oppositely directed unit rotations do likewise. This eliminates all aspects of the coupling of the motion to the reference system other than the reference point, reducing the particles to radiation, and bringing them to rest in the natural reference system. As seen in the spatial reference system, they become two photons moving outward in opposite directions from the point in the reference system at which the neutralization took place.

This neutralization, or “annihilation,” process becomes more difficult to accomplish as the particles increase in size and complexity, and takes place on a significant scale only in the sub-atomic range. However, full units of the magnetic rotation of the atom—units of inward rotational speed displacement—can be neutralized by combination with outward displacements of equal magnitude. The outward motions available for this purpose are ionization and thermal motion. When the total displacement of these motions reaches equality with that of a full unit of the magnetic rotation of an atom, or any full unit of that rotation, the existence of the rotational unit terminates, and its speed displacement reverts to the linear basis (radiation or kinetic energy).

As we saw earlier, the thermal ionization level is related to the temperature. The total outward speed displacement at which neutralization occurs is therefore reached at a specific temperature, a destructive temperature limit. Full ionization is attained at a level far below this limiting temperature. Inasmuch as it is the total outward displacement that enters into the neutralization process, rather than the thermal motion alone, the temperature of the destructive limit of an element depends on its atomic number. The heavier elements have more displacement in the form of ionization when all are fully ionized, and these elements therefore reach the same total displacement at lower temperatures.

When the temperature of an aggregate arrives at the destructive limit of the heaviest element present, this element reduces to one with less magnetic (two-dimensional) rotation, the difference in mass, t^3/s^3 , being converted to its one-dimensional equivalent, energy, t/s . As the rise in temperature continues, one after another of the elements meets the same fate in the order of decreasing atomic number.

CHAPTER 18

The Retreat from Reality

In the eight chapters from 9 to 17 (excluding Chapter 12) we have described the general features of electricity—both current electricity and electric charges—as they emerge from a development of the consequences of the postulates of the theory of the universe of motion. This development arrives at a picture of the place of electricity in the physical universe that is totally different from the one that we get from conventional physical theory. However, the new view agrees with the electrical observations and measurements, and is entirely consistent with empirical knowledge in related areas, whereas conventional theory is deficient in both respects. Thus there is ample justification for concluding that the currently accepted theories dealing with electricity are, to a significant degree, wrong.

This finding that an entire subdivision of accepted physical theory is not valid is difficult for most scientists to accept, particularly in view of the remarkable progress that has been made in the application of existing theory to practical problems. But neither a long period of acceptance nor a record of usefulness is sufficient to verify a theory. The history of science is full of theories that enjoyed general acceptance for long periods of time, and contributed significantly to the advance of knowledge, yet eventually had to be discarded because of fatal defects. Present-day electrical theory is not unique in this respect; it is just another addition to the long list of temporary solutions to physical problems.

The question then arises, How is it possible for errors of this magnitude to make their way into the accepted structure of physical theory? It is not difficult to find the answer. Actually, there are so many factors tending to facilitate acceptance of erroneous theories, and to resist parting with them after they are once accepted, that it has been something of an achievement to keep the error content of physical

theory as low as it is. The fundamental problem is that physical science deals with so many entities and phenomena whose basic nature is not understood. For example, present-day physics has no understanding of the nature of the electric charge. We are simply told that we must not ask; that the existence of charges has to be accepted as one of the given features of nature. This frees theory construction from the constraints that would normally apply. In the absence of an adequate understanding, it is possible to construct and secure acceptance of theories in which charges are assigned functions that are clearly seen to be incompatible with the place of electric charge in the pattern of physical activity, once that place is specifically defined.

None of the other basic entities of the physical universe—about six or eight of them, the exact number depending on the way in which the structure of fundamental theory is erected—is much, if any, better known than the electric charge. The nature of time, for instance, is even more of a mystery. But these entities are the foundation stones of physics, and in order to construct a physical theory it is necessary to make some assumptions about each of them. This means that present-day physical theory is based on some thirty or forty assumptions about entities that are almost totally unknown.

Obviously, the probability that *all* of these assumptions about the unknown are valid is near zero. Thus it is practically certain, simply from a consideration of the nature of its foundations, that the accepted structure of theory contains some serious errors.

In addition to the effects of the lack of understanding of the fundamental entities of the physical universe, there are some further reasons for the continued existence of errors in conventional physical theory that have their origin in the attitudes of scientists toward their subject matter. There is a general tendency, for instance, to regard a theory as firmly established if, according to the prevailing scientific opinion, it is the best theory of the subject that is currently available. As expressed by Henry Margenau, the modern scientist does not speak of a theory as true or false, but as “correct or incorrect relative to a given state of scientific knowledge.”⁶⁴

One of the results of this policy is that conclusions as to the validity of theories along the outer boundaries of scientific knowledge are

customarily reached without any consideration of the cumulative effect of the weak links in the chains of deductions leading to the premises of these theories. For example, we frequently encounter statements similar to the following:

The laws of modern physics virtually demand that black holes exist.⁶⁵

No one who accepts general relativity has found any way to escape the prediction that black holes must exist in our galaxy.⁶⁶

These statements tacitly assume that the reader accepts the “laws of modern physics” and the assertions of general relativity as incontestable, and that all that is necessary to confirm a conclusion—even a preposterous conclusion such as the existence of black holes—is to verify the logical validity of the deductions from these presumably established premises. The truth is, however, that the black hole hypothesis stands at the end of a long line of successive conclusions, included in which are more than two dozen pure assumptions. When this line of theoretical development is examined as a whole, rather than merely looking at the last step on a long road, it can be seen that arrival at the black hole conclusion is a clear indication that the line of thought has taken a wrong turn somewhere, and has diverged from physical reality. It will therefore be appropriate, in the present connection, to undertake an examination of this line of theoretical development, which originated with some speculations as to the nature of electricity.

The age of electricity began with a series of experimental discoveries: first, static electricity, positive* and negative*, then current electricity, and later the identification of the electron as the carrier of the electric current. Two major issues confronted the early theorists, (1) Are static and current electricity different entities, or merely two different forms of the same thing?, and (2) Is the electron only a charge, or is it a charged particle? Unfortunately, the consensus reached on question (1) by the scientific community was wrong. The theory of electricity thus took a wrong direction almost from the start. There was spirited opposition to this erroneous conclusion in the early days of electrical research, but Rowland’s experiment, in which he demonstrated that a moving charge has the magnetic

properties of an electric current, silenced most of the critics of the “one electricity” hypothesis.

The issue as to the existence of a carrier of electric charge—a “bare” electron—has not been settled in this manner. Rather, there has been a sort of a compromise. It is now generally conceded that the charge is not a completely independent entity. As expressed by Richard Feynman, “there is still ‘something’ there when the charge is removed.”⁶⁷ But the wrong decision on question (1) prevents recognition of the functions of the uncharged electron, leaving it as a vague “something” not credited with any physical properties, or any effect on the activities in which the electron participates. The results of this lack of recognition of the physical status of the uncharged electron, which we have now identified as a unit of electric quantity, were described in the preceding pages, and do not need to be repeated. What we will now undertake to do is to trace the path of a more serious retreat from reality that affects a large segment of present-day physical theory, and accounts for a major part of the difference between current theory and the conclusions derived from the postulates that define the universe of motion.

This theoretical development that we propose to examine originated as a result of the discovery of radioactivity and the identification of the three kinds of emanations from the radioactive substances as positively* charged alpha particles (helium atoms), negatively* charged electrons, and electromagnetic radiation. It was taken for granted that when certain particles are ejected from an atom during radioactivity, these particles must have existed in the atom prior to the radioactive disintegration. This conclusion does not seem so obvious today, when the photon of radiation (which no one suggests as a constituent of the undisturbed atom) is recognized as a particle, and a whole assortment of strange particles is observed to be emitted from atoms during high energy disintegrations. At any rate, it is clearly nothing more than an assumption.

An extension of this assumption led to the conclusion that the atom is a composite structure in which the emitted particles are the constituent parts. Some early suggestions as to the arrangement of the parts gained little support, but a discovery, in Rutherford’s

laboratory, that the mass of the atom is concentrated in a very small volume in the center of the space that it presumably occupies, led to the construction of the Rutherford atom-model, the prototype of the atom of modern physics. In this model the atom is viewed as a miniature analog of the solar system, in which negatively* charged electrons are in orbit around a positively* charged “nucleus.”

The objective of this present discussion is to identify the path that the development of theory on the basis of this atom-model has taken, and to demonstrate the fact that currently accepted theory along the outer boundaries of scientific knowledge, such as the theory that leads to the existence of black holes, rests on an almost incredible succession of pure assumptions, each of which has a finite probability—in some cases a very strong probability—of being wrong. As an aid in emphasizing the overabundance of these assumptions, we will number those that we identify as being definitely in the direct line of the theoretical development that leads eventually to the concepts of the black hole and the singularity.

In the construction of his model, Rutherford accepted the then prevailing concepts of the properties of electricity, including the two assumptions previously mentioned, and retained the assumption that the atom is constructed of separable parts. The first of the assumptions that he added will therefore be given the number 4. These new assumptions are:

- (4) The atom is constructed of positively* and negatively* charged components.
- (5) The positive* component, containing most of the mass, is located in a small nucleus.
- (6) Negatively* charged electrons are in orbit around the nucleus.
- (7) The force of attraction between unlike charges applied to motion of the electrons results in a stable orbital equilibrium.

This model met with immediate favor in scientific circles, but it was faced with two serious problems. The first was that the known behavior of unlike charges does not permit their coexistence at the very short distances in the atom. Even at substantially greater distances they neutralize each other. Strangely enough, little attention

was paid to this very important point. It was tacitly assumed (8) that the observed behavior of charges does not apply in this case, and that the hypothetical charges inside the atom are stable. There is no evidence whatever to support this assumption, but neither is there any evidence to contradict it, as the inside of the atom is unobservable. Here, as in many other areas of present-day physical theory, we are being asked to accept *absence of disproof* as the equivalent of proof.

Another of the problems encountered by the new theory involved the stability of the assumed electronic orbits. Here there was a direct conflict with empirical knowledge. From experiment it is found that charged objects moving in circular orbits (and therefore accelerated) lose energy and spiral in toward the center of the circle. On this basis the assumed electronic orbits would be unstable. This conflict was taken more seriously than the other, and remained a source of theoretical difficulty until Bohr “solved” the problem with another assumption, postulating, entirely *ad hoc*, that the constituents of the atom do not follow normal physical laws. He assumed (9) that the hypothetical electronic orbits are quantized, and can take only certain specific values, thus eliminating the spiraling effect.

At this point, further impetus was given to the development of the atom-model by the discovery of a positively* charged particle of mass one on the atomic weight scale. This particle, called the proton, was promptly assumed (10) to be the bare nucleus of the hydrogen atom. This led to the further assumption (11) that the nuclei of other atoms were made up of a varying number of protons. But here, again, there was a conflict with observation. According to the observed behavior of charged particles, the protons in the hypothetical nucleus would repel each other, and the nucleus would disintegrate. Again an *ad hoc* assumption was devised to rescue the atom-model. It was assumed (12) that an inward-directed “nuclear force” (of unknown origin) operates against the outward force of repulsion, and holds the protons in contact.

This assumed proton-electron composition quickly encountered difficulties, one of the most immediate being that in order to account for the various atoms and isotopes it had to be assumed that some of the electrons are located in the nucleus—admittedly a rather

improbable hypothesis. The theorists were therefore much relieved when a neutral particle, the neutron, was discovered. This enabled changing the assumed atomic composition to identify the nucleus as a combination of protons and neutrons (assumption 13). But the observed neutron is unstable, with an average life of only about 15 minutes. It therefore does not qualify as a possible constituent of a stable atom. So once more an *ad hoc* assumption was called upon. It was assumed (14) that the ordinarily unstable neutron becomes stable when it enters the atomic structure (where, fortunately for the hypothesis, it is undetectable if it exists).

As a result of the critical study to which the Bohr atom-model was subjected in the next few decades, this model, in its original form, was found untenable. Various “interpretations” of the model have therefore been offered as remedies for the defects in this original version. Each of these adds some further assumptions to those included in Bohr’s formulation, but none of these additions can be considered definitely in the main line of the theoretical development that we are following, and they will not be taken into account in the present connection. It should be noted, however, that all 14 of the assumptions that we have identified in the foregoing paragraphs enter into the theoretical framework of each modification of the atom-model. Thus all 14 are included in the premises of the “atom of modern physics,” regardless of the particular interpretation that is accepted.

It should also be noted that four of these 14 assumptions (numbers 8, 9, 12, and 14) have a status that is quite different from that of the others. These are *ad hoc assumptions*, untestable assumptions that are made purely for the purpose of evading conflicts with observation or firmly established theory. Assumption 12, which asserts the existence of a “nuclear force,” is a good example. There is no independent evidence that this assumed force actually exists. The only reason for assuming its existence is that the nuclear atom cannot survive without it. As one current physics textbook explains, “A very strong attractive force is needed to hold the nucleons in the nucleus.”⁶⁸ What the physicists are doing here is giving us an untestable excuse for the failure of the nuclear theory to pass the test of agreement with experience. Such evasive tactics are not new. In Aristotle’s physical system, which was the orthodox view of the

universe for nearly two thousand years, it was assumed that the planets were attached to transparent spheres that rotated around the earth. But according to the laws of motion, as they were understood at that time, this motion could not be maintained except by continual application of a force. So Aristotle employed the same device that his modern successors are using: the *ad hoc* assumption. He postulated the existence of angels who pushed the planets along in their respective orbits. The “nuclear force” of modern physics is the exact equivalent of Aristotle’s “angels” in all but language.

With the benefit of the additional knowledge that has been accumulated in the meantime, we of the present era have no difficulty in arriving at an adverse judgment on Aristotle’s assumption. But we need to recognize that this is an illustration of a general proposition. The probability that an untestable assumption about a physical entity or phenomenon is a true representation of physical reality is always low. This is an unavoidable consequence of the great diversity of physical existence. When one of these untestable assumptions is used in the *ad hoc* manner—that is, to evade a discrepancy or conflict—the probability that the assumption is valid is much lower.

All of these points are relevant to the question as to whether the present-day nuclear atom-model is a representation of physical reality. We have identified 14 assumptions that are directly involved in the main line of theoretical development leading to this model. These assumptions are sequential; that is, each adds to the assumptions previously made. It follows that unless *every one* of them is valid, the atom-model in its present form is untenable. The issue thus reduces to the question: What is the probability that *all* of these 14 assumptions are physically correct?

Here we need to consider the status of assumptions in the structure of scientific theory. The construction of physical theory is a process of applying reasoning to premises derived initially from experience. Where the application involves going from the general to the particular, the process is *deductive* reasoning, which is a relatively straightforward operation. To go from the particular to the general requires the application of *inductive* reasoning. This is a two-step process. First, a hypothesis is formulated by any one of a number of

means. Then the *hypothesis* is tested by developing its consequences and comparing them with empirical knowledge. Positive verification is difficult because of the great complexity of physical existence. It should be noted, in this connection, that agreement of the hypothesis with the observation that it was designed to fit does not constitute a verification. The hypothesis, or its consequences, must be shown to agree with *other* factual knowledge.

Because of the verification difficulties, it has been found necessary to make use, at least temporarily, of many hypotheses whose verification is incomplete. However, a prominent feature of “modern physics” is the extent to which the structure of theory rests on hypotheses that are entirely untested, and, in many cases, untestable. Hypotheses that are accepted and utilized without verification are *assumptions*. The use of assumptions is a legitimate feature of theory or model construction. But in view of the substantial uncertainty as to their validity that always exists, the standard scientific practice is to avoid pyramiding them. One or two steps into the unknown are considered to be in order, but some consolidation of the exposed positions is normally regarded as essential before a further unsupported advance is undertaken.

The reason for this can easily be seen if we consider the way in which the probability of validity is affected. Because of the complexity of physical existence mentioned earlier, the probability that an untestable assumption is valid is inherently low. In each case, there are *many* possibilities to be conceived and taken into account. If each assumption of this kind has an even chance (50 percent) of being valid, there is some justification for using one such assumption in a theory, at least tentatively. If a second untestable assumption is introduced, the probability that both are valid becomes one in four, and the use of these assumptions as a basis for further extension of theory is a highly questionable practice. If a third such assumption is added, the probability of validity is only one in eight, which explains why pyramiding assumptions is regarded as unsound.

A consideration of the points brought out in the foregoing paragraphs puts the status of the nuclear theory into the proper perspective. The 14 steps in the dark that we have identified in the path of development

of the currently accepted atom-model are totally unprecedented in physical science. The following comment by Abraham Pais is appropriate:

Despite much progress, Einstein's earlier complaint remains valid to this day. "The theories which have gradually been associated with what has been observed have led to an unbearable accumulation of individual assumptions."⁶⁹

Of course, it is possible for an assumption to be upgraded to the status of established knowledge by discovery of confirmatory evidence. This is what happened to the assumption as to the existence of atoms. But none of the 14 numbered assumptions identified in the preceding discussion has been similarly raised to a factual status. Indeed, some of them have lost ground over the years. For example, as noted earlier, the assumption that emission of certain particles from an atom during a decay process indicates that these particles existed in the atom before decay, assumption (3), has been seriously weakened by the large increase in the number of new particles that are being emitted from atoms during high energy processes. The present uncritical acceptance of the nuclear atom-model is not a result of more empirical support, but of increasing familiarity, together with the absence (until now) of plausible alternatives. A comment by N. R. Hanson on the quantum theory, one of the derivatives of the nuclear atom model, is equally applicable to the model itself. This theory, he says, is "conceptually imperfect" and "riddled with inconsistencies." Nevertheless, it is accepted in current practice because "it is the only extant theory capable of dealing seriously with microphenomena."⁷⁰

The existence, or non-existence, of alternatives has no bearing, however, on the question we are now examining, the question as to whether the nuclear atom-model is a true representation of physical reality. Neither general acceptance nor long years of freedom from competition has any bearing on the validity of the model. Its probability of being correct depends on the probability that the 14 assumptions on which it rests are *all* individually valid. Even if no ad hoc assumptions were involved, this composite probability, the product of the individual probabilities, would be low because of the cumulative effect. This line

of theoretical development is the kind of product that Einstein called “an unbearable accumulation of individual assumptions.” Even if we assume the relatively high value of 90 percent for the probability of the validity of each individual assumption, the probability that the final result, the atom-model, is correct would be less than one in four. When the very low probability of the four purely ad hoc assumptions is taken into account, it is evident that the probability of the nuclear atom-model, “the atom of modern physics,” being a correct representation of physical reality is close to zero.

This conclusion derived from an examination of the foundations of the currently accepted model will no doubt be resisted—and probably resented—by those who are accustomed to the confident assertions in the scientific literature. But it is exactly what many of those who played leading roles in the development of the long list of assumptions leading to the present version of the nuclear theory have been telling us. These scientists know that the construction of the model in terms of electrons moving in orbits around a positively* charged nucleus does not mean that such entities actually exist in an atom, or behave in the manner specified in the theory. Erwin Schrödinger, for instance, emphasized that the model is “only a mental help, a tool of thought.”⁷¹ and asserted that if the question, “Do the electrons actually exist in these orbits within the atom?” is asked, it “is to be answered with a decisive *No*.”⁷² Werner Heisenberg, another of the architects of the modern version of Bohr’s atom-model, tells us that the physicists’ atom does not even “exist objectively in the same sense as stones or trees exist.”⁷³ It is, “in a way, only a symbol,”⁷⁴ he says.

These statements, applying specifically to the nuclear theory of the atom, that have been made by individuals who know the true status of the assumptions that entered into the construction of that theory, agree with the conclusions that we have reached on the basis of probability considerations. Thus the confident statements that appear throughout the scientific literature, asserting that the nature of the atomic structure is now “known,” are wholly unwarranted. A hypothesis that is “only a mental help” is not a representation of reality. A theoretical line of development that culminates in nothing more than a “symbol” or a “tool of thought” is not an exploration of the real world; it is an excursion into the land of fantasy.

The finding that the nuclear atom-model rests on false premises does not necessarily invalidate the currently accepted *mathematical* relationships derived from it, or suggested by it. This may appear contradictory, as it implies that a wrong theory may lead to correct answers. However, the truth is that the conceptual and mathematical aspects of physical theories are, to a large extent, independent. As Feynman puts it, "Every theoretical physicist who is any good knows six or seven different theoretical representations for exactly the same physics."⁷⁴ Such a physicist recognizes this many different conceptual explanations that agree with the mathematical relations. A major reason for this is that the mathematical relations are usually identified first, and an explanation in the form of a theory is developed later as an interpretation of the mathematics. As noted earlier, many such explanations are almost always possible in each case. In the course of the investigation on which this present work is based, this has been found to be true even where the architects of present-day theory contend that "there is no other way."

Since the practical applications of a theory are primarily mathematical, or quantitative, one might be led to ask, Why do we want an explanation? Why not just use the mathematics without any concern as to their meaning? The answer is that while the established mathematical relations may serve the specific purposes for which they were developed, they cannot be safely extrapolated beyond the ranges of conditions over which they have been tested, and they make no contribution toward an understanding of relations in other areas. On the contrary, they lead to wrong conclusions, and constitute roadblocks in the way of identifying the correct principles and relations in related areas.

This is what has happened as a result of the assumptions that were made in the course of developing the nuclear atom-model. Once it was assumed that the atom is composed primarily of oppositely charged particles, and some valid mathematical relations were developed and expressed in terms of this concept, the prevailing tendency to accept mathematical agreement as proof of validity, together with the absence (until now) of any serious competition, elevated this product of multiple assumptions to the level of an accepted fact. "Today we know that the atom consists of a positively charged

nucleus composed of protons and neutrons surrounded by negatively charged electrons.” This positive statement, or its equivalent, can be found in almost every physics textbook. But any proposition that rests on assumptions is hypothesis, not knowledge. Classifying a model that rests upon more than a dozen independent assumptions, mostly untestable, and including several of the inherently dubious “ad hoc” variety, as “knowledge” is a travesty on science.

When the true status of the nuclear atom-model is thus identified, it should be no surprise to find that the development of the theory of the universe of motion reveals that the atom actually has a totally different structure. We now find that it is not composed of individual particles, and in its normal state it contains no electric charges. This new view of atomic structure was derived by deduction from the postulates that define the universe of motion, and it therefore participates in the verification of the Reciprocal System of theory as a whole. However, in view of the crucial position of the nuclear theory in conventional physics it is advisable to make it clear that this currently accepted theory is almost certainly wrong, *on the basis of current physical knowledge*, even without the additional evidence supplied by the present investigation, and that some of the physicists who were most active in the construction of the modern versions of the nuclear model concede that it is not a true representation of physical reality. This is the primary purpose of the present chapter.

In line with this objective, the most significant of the errors introduced into electric and magnetic theory by acceptance of this erroneous model of atomic structure have been identified in the preceding pages. But this is not the whole story. This product of “an unbearable accumulation of individual assumptions” has had an even more detrimental effect on astronomy. The errors that it has introduced into astronomical thought will be discussed in detail in Volume III, but it will be appropriate at this time to point out why astronomy has been particularly vulnerable to an erroneous assumption of this nature.

The magnitudes of the basic physical properties extend through a much wider range in the astronomical field than in the terrestrial environment. A question of great significance, therefore, in the

study of astronomical phenomena, is whether the physical laws and principles that apply under terrestrial conditions are also applicable under the extreme conditions to which many astronomical objects are subjected. Most scientists are convinced, largely on philosophical, rather than scientific, grounds, that the same physical laws do apply throughout the universe. The results obtained by development of the consequences of the postulates that define the universe of motion agree with this philosophical assumption. However, there is a general tendency to interpret this principle of universality of physical law as meaning that *the laws that have been established as applicable to terrestrial conditions are applicable throughout the universe*. This is something entirely different, and our findings do not support it.

The error in this interpretation of the principle stems from the fact that most physical laws are valid, in the form in which they are usually expressed, only within certain limits. Many of the currently accepted laws applicable to solids, for example, do not apply at temperatures above the melting points of the various material substances. The prevailing interpretation of the uniformity principle carries with it the unstated assumption that there are no such limits applicable to the currently accepted laws and principles other than those that are recognized in present-day practice. In view of the very narrow range of conditions through which these laws and principles have been tested, this assumption is clearly unjustified, and our findings now show that it is definitely incorrect. We find that while it is true that the same laws and principles are applicable throughout the universe, most of the basic laws are subject to certain modifications at critical magnitudes, which often exceed the limiting magnitudes experienced on earth, and are therefore unknown to present-day science. Unless a law is so stated that it provides for the existence and effects of these critical magnitudes, it is *not* applicable to the universe as a whole, however accurate it may be within the narrow terrestrial range of conditions.

One property of matter that is subject to an unrecognized critical magnitude of this nature is density. In the absence of thermal motion, each type of material substance in the terrestrial environment has a density somewhere in the range from 0.075 (hydrogen) to 22.5 (osmium and iridium), relative to liquid water at 4 °C as 1.00. The

average density of the earth is 5.5. Gases and liquids at lower densities can be compressed to this density range by application of sufficient pressure. Additional pressure then accomplishes some further increase in density, but the increase is relatively small, and has a decreasing trend as the pressure rises. Even at the pressures of several million atmospheres reached in shock wave experiments, the density was only increased by a factor of about two. Thus the maximum density to which the contents of the earth could be raised by application of pressure is not more than about 15.

The density of most of the stars of the white dwarf class is between 100,000 and 1,000,000. There is no known way of getting from a density of 15 to a density of 100,000. And present-day physics has no general theory from which an answer to this problem can be deduced. So the physicists, already far from the solid ground of reality with their hypotheses based on an atom-model that is “only a symbol,” plunge still farther into the realm of the imagination by adding more assumptions to the sequence of 14 included in the nuclear atom-model. It is first assumed (15) that at some extremely high pressure the hypothetical nuclear structure collapses, and its constituents are compressed into one mass, eliminating the vacant space in the original structure, and increasing the density to the white dwarf range.

How the pressure that is required to produce the “collapse” is generated has never been explained. The astronomers generally assume that this pressure is produced at the time when, according to another assumption (16), the star exhausts its fuel supply.

With its fuel gone it [the star] can no longer generate the pressure needed to maintain itself against the crushing force of gravity.⁷⁵

But fluid pressure is effective in all directions; down as well as up. If the “crushing force of gravity” is exerted against a gas rather than directly against the central atoms of the star, it is transmitted undiminished to those atoms. It follows that the pressure against the atoms is not altered by a change of physical state due to a decrease in temperature, except to the extent that the dimensions of the star may be altered. When it is realized that the contents of ordinary stars, those of the main sequence, are already in a condensed state (a point

discussed in detail in Volume III), it is evident that the change in dimensions is too small to be significant in this connection. The origin of the hypothetical “crushing pressure” thus remains unexplained.

Having assumed the fuel supply exhausted, and the star cooled down, in order to produce the collapse, the theorists find it necessary to reheat the star, since the white dwarfs are relatively hot, as their name implies, rather than cold. So again they call upon their imaginations and come up with a new assumption to take care of the problem. They assume (17) that when the atomic structure collapses, the matter of the star enters a new state. It becomes *degenerate matter*, and acquires a new set of properties, among which is the ability to generate a new supply of energy to account for the observed temperature of the white dwarf stars.

Even with the wide latitude for further assumptions that is available in this purely imaginary situation, the white dwarf hypothesis could not be extended sufficiently to encompass all of the observed types of extremely dense stars. To meet this problem it was assumed (18) that the collapse which produces the white dwarf is limited to relatively small stars, so that the white dwarfs do not exceed a limiting mass of about two solar masses. Larger stars are assumed (19) to explode rather than merely collapse, and it is further assumed (20) that the pressure generated by such an explosion is sufficient to compress the residual matter to such a degree that the hypothetical constituents of the atoms are all converted into neutrons, producing a neutron star (currently identified with the observed objects known as pulsars). There is no evidence to support this assumption. The existence of a process that accomplishes such a conversion under pressure is itself an assumption (21), and the concept of a neutron star requires the further assumption (22) that neutrons can exist as stable independent particles under the assumed conditions.

Although this is the currently orthodox explanation of the origin of the pulsars, it is viewed rather dubiously even by some of the astronomers. Martin Harwit, for instance, concedes that “we have no theories that satisfactorily explain just how a massive star collapses to become a neutron star.”⁷⁶

The neutron star, too, is assumed to have a limiting mass. It is assumed (23) that the compression due to the more powerful explosion

of the larger star reduces the volume of the residual aggregate enough to enable its self-gravitation to continue the compression. It is then further assumed (24) that the reduction of the size of the aggregate eventually reaches the point where the gravitational force is so great that radiation cannot escape. What then exists is a black hole.

While it is not generally recognized as such, the “self-gravitation” concept, in application to atoms, is another assumption (25). Observations show only that gravitation operates *between* atoms or independent particles. The hypothesis that it is also applicable within atoms is derived from Einstein’s general theory of relativity, but since there is no *proof* of this theory (the points that have thus far been adduced in its favor are merely *evidence*) this derivation does not alter the fact that the hypothesis of gravitation within atoms rests on an assumption.

Most astronomers who accept the existence of black holes apparently prefer to look upon these objects as the limiting state of physical existence, but others recognize that if self-gravitation is a reality, and if it is once initiated, there is nothing to stop it at an intermediate stage such as the black hole. These individuals therefore assume (26) that the contraction process continues until the material aggregate is reduced to a mere point, a *singularity*.

This line of thought that we have followed from the physicists’ concept of the nature of electricity to the nuclear model of atomic structure, and from there to the singularity, is a good example of the way in which unrestrained application of imagination and assumption in theory construction leads to ever-increasing levels of absurdity—in this case, from atomic “collapse” to degenerate matter to neutron star to black hole to singularity. Such a demonstration that extension of a line of thought leads to an absurdity, the *reductio ad absurdum*, as it is called, is a recognized logical method of disproving the validity of the premises of that line of thought. The physicist who tells us that “the laws of modern physics virtually demand that black holes exist” is, in effect, telling us that there is something wrong with the laws of modern physics. In the preceding pages we have shown just what is wrong: too much of the foundation of conventional physical theory rests on untestable assumptions and “models.”

The physical theory derived by development of the consequences of the postulates that define the universe of motion differs very radically from current thought in some areas, such as astronomy, electricity, and magnetism. Many scientists find it hard to believe that the investigators who constructed the currently accepted theories could have made so many mistakes. It should be emphasized, therefore, that the profusion of conflicts between present-day ideas and our findings does not indicate that the previous investigators have made a *multitude of errors*. What has happened is that they have made *a few serious errors* that have had a *multitude of consequences*.

The astronomical theories based on the nuclear atom-model that have been mentioned in this chapter provide a good example of how one basic error distorts the pattern of thinking over a wide area. In this case, an erroneous theory of the structure of the atom leads to an erroneous theory of extremely high density, which then results in the construction of erroneous theories of all of the astronomical objects composed of ultradense matter; not only the white dwarfs, but also quasars, pulsars, x-ray emitters, and compact galactic cores. Once the pyramiding of assumptions begins, such spurious results are inevitable.

CHAPTER 19

Magnetostatics

As we saw in the preceding pages, one of the principal obstacles to the development of a more complete and consistent theory of electrical phenomena has been the exaggerated significance that has been attached to the points of similarity between static and current electricity, an attitude that has fostered the erroneous belief that only one entity, electric charge, is involved in the two types of phenomena. The same kind of a mistake has been made in a more complete and categorical manner in the current view of magnetism. While insisting that electrostatic and current phenomena are simply two aspects of the same thing, contemporary scientific opinion concedes that there is enough difference between the two to justify a separate category of electrostatics for the theoretical aspects of the static phenomena. But if *magnetostatics*, the corresponding branch of magnetism, is mentioned at all in modern physics texts, it is usually brushed off as an “older approach” that is now out of date. Strictly static concepts, such as that of magnetic poles, are more often than not introduced somewhat apologetically.

The separation of physical fields of study into more and more subdivisions has been a feature of scientific activity throughout its history. Here in the magnetostatic situation we have an example of the reverse process, a case in which a major subdivision of physics has succumbed to cannibalism. Magnetostatics has been swallowed by a related, but quite different, phenomenon, *electromagnetism*, which will be considered in Chapter 21. There are many similarities between the two types of magnetic phenomena, just as there are between the two kinds of electricity. In fact, the quantities in terms of which the magnetostatic magnitudes are expressed are defined mainly by electromagnetic relations. But this is not by any means sufficient to justify the current belief that only one entity is involved. The sub-

ordinate status that conventional physics gives to purely magnetic phenomena is illustrated by the following comment from K. W. Ford:

As theoretical physicists see it, magnetism in our world is merely an accidental by-product of electricity; it exists only as a result of the motion of electrically charged particles.⁷⁷

The implication of a confident statement of this kind is that the assertions which it makes are reasonably well established. In fact, however, this assertion that magnetism exists only as a result of the motion of electrically charged particles is based entirely on unsubstantiated assumptions. The true situation is more accurately described by the following quotation from a physics textbook:

It is only within the past thirty years or so that models tying together these two sources of magnetism [magnets and electro-magnetism] have been developed. These models are far from perfect, even today, but at least they have convinced most people that there is really only one source of magnetic fields: all magnetic fields come from moving electric charges.⁷⁸

What this text is, in effect, telling us is that the idea does not work out very well in practice, but that it has been accepted by majority vote anyway. A prominent American astronomer, J. N. Bahcall, has pointed out that “We frequently settle important scientific issues by acclamation rather than observation.”⁷⁹ The uncritical acceptance of the “far from perfect” models of magnetism is a good example of this unscientific practice.

A strange feature of the existing situation is that after having come to this conclusion that magnetism is merely a by-product of electricity, one of the ongoing activities of the physicists is a search for the magnetic analog of the mobile electric charge, the electron. Again quoting K. W. Ford:

An electric particle gives rise to an electric field, and when it moves it produces a magnetic field as a secondary effect. For symmetry’s sake there should be magnetic particles that give rise to magnetic fields and in motion produce electric fields in the same way that moving electric particles create magnetic fields.

This author admits that “So far the magnetic monopole has frustrated all its investigators. The experimenters have failed to find

any sign of the particle.” Yet this will-o’-the-wisp continues to be pursued with an ardor that invites caustic comments such as this:

It is remarkable how the lack of experimental evidence for the existence of magnetic monopoles does not diminish the zeal with which they are sought.⁸⁰

Ford’s contention is that “the apparent non-existence of monopole particles now presents physicists with a paradox that they cannot drop until they have found an explanation.” But he (unintentionally) supplies the answer to the paradox when he closes his discussion of the monopole situation with this statement:

What concerns the physicist is that, in defiance of symmetry and all the known laws, no magnetic particle so far has been created or found anywhere.

Whenever the observed facts “defy” the “known laws” and the current understanding of the application of symmetry relations to any given situation, it can safely be concluded that the current understanding of symmetry and at least some of the “known laws” are wrong. In the present case, any critical appraisal will quickly show not only that a number of the premises from which the conclusion as to the existence of magnetic monopoles is derived are pure assumptions without factual support, but also that there is a definite contradiction between two of the key assumptions.

As explained by Ford, the magnetic monopole for which the physicists are searching so assiduously is a particle which “gives rise to magnetic fields”; that is, a magnetic charge. If such a particle existed, it would, of course, exhibit magnetic effects due to the charge. But this is a direct contradiction of the prevailing assumption that magnetism is a “by-product of electricity.” The physicists cannot have it both ways. If magnetism is a by-product of electricity (that is, electric charges, in current thought), then there cannot be a magnetic charge, a source of magnetic effects, analogous to the electric charge, a source of electric effects. On the other hand, if a particle with a magnetic charge (a magnetic monopole) does exist, then the physicists’ basic theory of magnetism, which attributes all magnetic effects to electric currents, is wrong.

It is obvious from the points brought out in the theoretical development in the preceding pages that the item of information which has been missing is an understanding of the physical nature of magnetism. As long as magnetism is assumed to be a by-product of electricity, and electricity is regarded as a given feature of nature, incapable of explanation, there is nothing to guide theory into the proper channels. But once it is recognized that magnetostatic phenomena are due to magnetic charges, and that such a charge is a type of motion—a rotational vibration—the situation is clarified almost automatically. Magnetic charges do, indeed, exist. Just as there are electric charges, which are one-dimensional rotational vibrations acting in opposition to one-dimensional rotations, there are magnetic charges, which are two-dimensional rotational vibrations acting in opposition to two-dimensional rotations. The phenomena due to charges of this nature are the subject matter of magnetostatics. Electromagnetism is a different phenomenon that is also two-dimensional, but involves motion of a continuous, rather than vibratory, nature.

The two-dimensionality is the key to understanding the magnetic relations, and the failure to recognize this basic feature of magnetism is one of the primary causes of the confusion that currently exists in many areas of magnetic theory. The two dimensions of the magnetic charge and electromagnetism are, of course, scalar dimensions. The motion components in the second dimension are not capable of direct representation in the conventional spatial reference system, but they have indirect effects that are observable, particularly on the effective magnitudes. Lack of recognition of the vibrational nature of electrostatic and magnetostatic motions, which distinguishes them sharply from the continuous motions involved in current electricity and electromagnetism, also contributes significantly to the confusion. Magnetostatics resembles electromagnetism in those respects in which the number of effective dimensions is the determining factor, whereas it resembles electrostatics in those respects in which the determining factor is the vibrational character of the motion.

Our findings show that the absence of magnetic monopoles is not a “defiance of symmetry.” The symmetry exists, but a better understanding of the nature of electricity and magnetism is required before it can be recognized. There is symmetry in the electric and

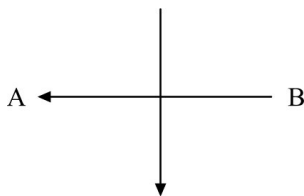
magnetic relations, and in some respects it is the kind of symmetry envisioned by Ford and his colleagues. One type of magnetic field is produced in the same manner as an electric field, just as Ford suggests in his explanation of the reasoning underlying the magnetic monopole hypothesis. But it is not an “electric particle” that produces an electric field; it is a certain kind of motion—a rotational vibration—and a magnetic field is produced by a similar rotational vibration. The electric current, a translational motion of a particle (the uncharged electron) in a conductor, produces a magnetic field. As we will see in Chapter 21, a translational motion of a magnetic field likewise produces an electric current in a conductor. Here, too, symmetry exists, but not the kind of symmetry that would call for a magnetic monopole.

The magnetic force equation, the expression for the force between two magnetic charges, is identical with the Coulomb equation, except for the factor $\frac{1}{s}$ introduced by the second scalar dimension of motion in the magnetic charge. The conventional form of the equation is $F = MM'/d^2$. As in the other primary force equations, the terms M' and d^2 are dimensionless. From the general principles applying to these force equations, as defined in Chapter 14, the missing term in the magnetic equation, analogous to $\frac{1}{s}$ in the Coulomb equation, is $\frac{1}{t}$. The space-time dimensions of the magnetic equation are then $F = t^2/s^2 \times \frac{1}{t} = \frac{1}{s^2}$.

Like the motion that constitutes the electric charge, and for the same reasons, the motion that constitutes the magnetic charge has the outward scalar direction. But since magnetic rotation is necessarily positive (time displacement) in the material sector, all stable magnetic charges in that sector have displacement in space (negative), and there is no independent magnetic phenomenon corresponding to the negative* electric charge. In this case there is no established usage that prevents applying the designations that are consistent with the rotational terminology, and we will therefore refer to the magnetic charge as negative, rather than using the positive* designation, as in application to the electric charge.

Although positive magnetic charges do not exist in the material environment, except under the influence of external forces in a situation that will be discussed later, the two-dimensional character of the magnetic charge introduces an orientation effect not present

in the electric phenomena. All one-dimensional (electric) charges are alike; they have no distinguishing characteristic whereby they can be subdivided into different types or classes. But a two-dimensional (magnetic) charge consists of a rotational vibration in the dimension of the reference system and another in a second scalar dimension independent of the first, and therefore perpendicular to it in a geometrical representation. The rotation with which this second rotational vibration is associated divides the atom into two halves that can be separately identified. On one side of this dividing line the rotation appears clockwise to observation. The scalar direction of the magnetic charge on this side is therefore outward from a clockwise rotation. A similar charge on the opposite side is a motion outward from a counterclockwise rotation.



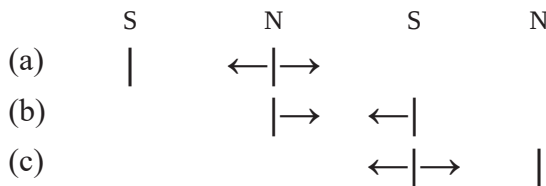
The above diagram shows the axes of rotation in one position, with the directions of rotation indicated by the arrowheads. Let us say that this represent a rotation that is clockwise as seen from point A. In this case the direction of the scalar resultant of the motion in the second dimension is toward point A. Now let us look at the same rotation from point B. Here the rotation is counterclockwise and the scalar resultant of the motion in the second dimension is directed away from the point of observation. Here, then, the clockwise rotation produces an attraction, whereas the counterclockwise rotation produces a repulsion. Thus the two-dimensional rotational vibration that constitutes the magnetic charge exerts an attractive force in one direction and a repulsive force in the opposite direction.

The unit of magnetic charge applies to only one of the two rotating systems of the atom. Each atom therefore acquires two charges, which occupy the positions described in the preceding paragraph, and are oppositely directed. Each atom of a magnetic, or magnetized, substance thus has two poles, or centers of magnetic effect at opposite ends of one of the axes of rotation. These are analogous

to the magnetic poles of the earth, and are named accordingly, as a north pole, or north-seeking pole, and a south pole.

These poles constitute scalar reference points, as defined in Chapter 12. The effective direction of the rotational vibration that constitutes the charge located at the north pole is outward from the north reference point, while the effective direction of the charge centered at the south pole is outward from the south reference point. The interaction of two magnetically charged atoms therefore follows the same pattern as the interaction of electric charges. As illustrated in Fig. 22, which is identical with Fig. 20, Chapter 13, except that it substitutes poles for charges, two north poles (line a) move outward from north reference points, and therefore outward from each other. Two south poles (line c) similarly move outward from each other. But, as shown in line b, a north pole moving outward from a north reference point is moving toward a south pole that is moving outward from a south reference point. Thus like poles repel and unlike poles attract.

FIGURE 22



On this basis, when two magnetically charged atoms are brought into proximity, the north pole of one atom is drawn to the south pole of the other. The resulting structure is a linear combination of a north pole, a neutral combination of two poles, and a south pole. Addition of a third magnetically charged atom converts this south pole to a neutral combination, but leaves a new south pole at the new end of the structure. Further additions of the same kind can then be made, limited only by thermal and other disruptive forces. A similar linear array of atoms with north and south poles at opposite ends can be produced by introducing atoms of magnetizable matter between the magnetically charged atoms of a two-atom combination. Separation of this structure at any point breaks a neutral combination, and leaves north and south poles at the ends of each segment. Thus no matter

how finely magnetic material is divided, there are always north and south poles in every piece of the material.

Because of the directional character of the magnetic forces they are subject to shielding in the same manner as electric forces. The gravitational force, on the other hand, cannot be screened off or modified in any way. Many observers have regarded this as an indication that the gravitational force must be something of a fundamentally different nature. This impression has been reinforced by the difficulty that has been experienced in finding an appropriate place for gravitation in basic physical theory. The principal objective of the theorists currently working on the problem of constructing a “unified theory,” or a “grand unified theory,” of physics is to find a place for gravitation in their theoretical structure.

Development of the theory of the universe of motion now shows that gravitation, static electricity, and magnetostatics are phenomena of the same kind, differing only in the number of effective scalar dimensions. Because of the symmetry of space and time in this universe, every kind of force (motion) has an oppositely directed counterpart. Gravitation is no exception. It takes place in time as well as in space, and is therefore subject to the same differentiation between positive and negative as that which we find in electric forces. But in the material sector of the universe the net gravitational effect is always in space—that is, there is no effective negative gravitation—whereas in the cosmic sector it is always in time. And since gravitation is three-dimensional, there cannot be any directional differentiation of the kind that we find in magnetism.

Because of the lack of understanding of the true relation between electromagnetic and gravitational phenomena, conventional physical science has been unable to formulate a theory that would apply to both. The approach that has been taken to the problem is to *assume* that electricity is fundamental, and to erect the structure of physical theory on this foundation, further assumptions being made along the way as required in order to bring the observations and measurements into line with the electrically based theory. Gravitation has thus been left with the status of an unexplained anomaly. This is wholly due to the manner in which the theories have been constructed, not to any peculiarity of gravitation. If the approach had been reversed, and

physical theory had been constructed on the basis of the assumption that gravitation is fundamental, electricity and magnetism would have been the “undigestable” items. The kind of a unified theory that the investigators have been trying to construct can only be attained by a development, such as the one reported in this work, that rests on a solid foundation of understanding in which each of these three basic phenomena has its proper place.

Aside from the effects of the difference in the number of scalar dimensions, the properties of the rotational vibration that constitutes a magnetic charge are the same as those of the rotational vibration that constitutes an electric charge. Magnetic charges can therefore be *induced* in appropriate materials. These materials in which magnetic charges are induced behave in the same manner as permanent magnets. In fact, some of them *become* permanent magnets when charges are induced in them. However, only a relatively small number of elements are capable of being magnetized to a significant degree; that is, have the property known as *ferromagnetism*.

The conventional theories of magnetism have no explanation of the restriction of magnetization (in this sense) to these elements. Indeed, these theories would seem to imply that it should be a general property of matter. On the basis of the assumptions previously mentioned, the electrons which conventional theory regards as constituents of atoms are miniature electromagnets, and produce magnetic fields. In most cases, it is asserted, the magnetic fields of these atoms are randomly oriented, and there is no net magnetic resultant. “However, there are a few elements in whose atoms the fields from the different electrons don’t exactly cancel, and these atoms have a net magnetic field... in a few materials... the magnetic fields of the atoms line up with each other.”⁸¹ Such materials, it is asserted, have magnetic properties. Just why these few elements should possess a property that most elements do not have is not specified.

For an explanation in terms of the theory of the universe of motion we need to consider the nature of the atomic motion. If a two-dimensional positive rotational vibration is added to the three-dimensional combination of motions that constitutes the atom it modifies the magnitudes of those motions, and the product is not the same atom with a magnetic charge; it is an atom of a different kind. The results of such additions

will be examined in Chapter 24. A magnetic charge, as a distinct entity, can exist only where an atom is so constituted that there is a portion of the atomic structure that can vibrate two-dimensionally independently of the main body of the atom. The requirements are met, so far as the magnetic rotation is concerned, where this rotation is asymmetric; that is, there are n displacement units in one of the two magnetic dimensions and $n+1$ in the other.

On this basis, the symmetrical B groups of elements, which have magnetic rotations 1-1, 2-2, 3-3, and 4-4, are excluded. While the magnetic charge has no third dimension, the electric rotation with which it is associated in the three-dimensional motion of the atom must be independent of that associated with the remainder of the atom. The electric rotational displacement must therefore exceed 7, so that one complete unit (7 displacement units plus an initial unit level) can stay with the main body of the magnetic rotation, while the excess applies to the magnetic charge. Furthermore, the electric displacement must be positive, as the reference system cannot accommodate two different negative displacements (motion in time) in the same atomic structure. The electronegative divisions (III and IV) are thus totally excluded. The effect of all of these exclusions is to confine the magnetic charges to Division II elements of Groups 3A and 4A.

In Group 3A the first element capable of taking a magnetic charge in its normal state is iron. This number one position is apparently favorable for magnetization, as iron is by far the most magnetic of the elements, but a theoretical explanation of this positional advantage is not yet available. The next two elements, cobalt and nickel, are also magnetic, as their electric displacement is normally positive. Under some special conditions, the displacements of chromium (6) and manganese (7) are increased to 8 and 9 respectively by reorientation relative to a new zero point, as explained in Chapter 18 of Volume I. These elements are then also able to accept magnetic charges.

According to the foregoing explanation of the atomic characteristics that are re-quired in order to permit acquisition of a magnetic charge, the only other magnetic (in this sense) elements are the members of Division II of Group 4A (magnetic displacements 4-3). This theoretical expectation is consistent with observation, but there are some, as yet unexplained, differences between the magnetic

behavior of these elements and that of the Group 3A elements. The magnetic strength is lower in the 4A group. Only one of the elements of this group, gadolinium, is magnetic at room temperature, and this element does not occupy the same position in the group as iron, the most magnetic element of Group 3A. However, samarium, which is in the iron position, does play an important part in many magnetic alloys. Gadolinium is two positions higher in the atomic series, which may indicate that it is subject to a modification similar to that applying to the lower 3A elements, but oppositely directed.

If we give vanadium credit for some magnetic properties on the strength of its behavior in some alloys, all of the Division II elements of both the 3A and 4A groups have a degree of magnetism under appropriate conditions. The larger number of magnetic elements in Group 4A is a reflection of the larger size of this 32 element group, which puts 12 elements into Division II. There are a number of peculiarities in the relation of the magnetic properties of these 4A elements, the rare earths, to the positions of the elements in the atomic series that are, as yet, unexplained. They are probably related to the other still unexplained irregularities in the behavior of these elements that were noted in the discussions of other physical properties. The magnetic capabilities of the Division II elements and alloys carry over into some compounds, but the simple compounds such as the binary chlorides, oxides, etc. tend to be non-magnetic; that is, incapable of accepting magnetic charges of the ferromagnetic type.

In undertaking an examination of individual magnetic phenomena, our first concern will be to establish the correct dimensions of the quantities with which we will be working. This is an operation that we have had to carry out in every field that we have investigated, but it is doubly important in the case of magnetism because of the dimensional confusion that admittedly exists in this area. The principal reason for this confusion is the lack in conventional physical theory of any valid general framework to which the dimensional assignments of electric and magnetic quantities can be referred. The customary assignment of dimensions on the basis of an analysis into mass, length, and time components produces satisfactory results in the mechanical system of quantities. Indeed, all that is necessary to convert these mechanical MLT assignments to the correct space-time dimensions is to recognize

the $t^{3/3}$ dimensions of mass. But extending this MLT system to electric and magnetic quantities meets with serious difficulties. Malcolm McCaig makes the following comment:

Very contradictory statements have been made about the dimensions of electrical quantities. While some writers state that it is impossible to express the dimensions of all electrical and magnetic quantities in terms of mass, length, and time, others such as Jeans and Nicolson do precisely that.⁸²

The nature of the problem that the theorists face in attempting to arrive at an accurate and consistent set of MLT dimensions can be seen by comparing the dimensions that have been assigned to one of the basic electric quantities, electric current, with the correct space-time dimensions that we have identified in the preceding pages. Charge (Q), in MLT terms, is asserted to have the dimensions $M^{1/2} L^{3/2} T^{-1}$. After converting mass into its space-time equivalent, this expression becomes $(t^{3/3})^{1/2} \times s^{3/2} \times t^{-1} = t^{1/2}$. The correct dimensions are $1/s$. The reason for the discrepancy is that the MLT dimensions are taken from the force equations, and therefore reflect the errors in the conventional interpretation of those equations. The further error due to the lack of distinction between electric charge and electric quantity is added when dimensions are assigned to the electric current. The MLT dimensions are $M^{1/2} L^{3/2} T^{-2}$. Again converting to space-time dimensions, we have $(t^{3/3})^{1/2} \times s^{3/2} \times t^{-2} = 1/t^{1/2}$. The correct dimensions are s/t .

The SI system and its immediate predecessors avoid a part of the problem by abandoning the effort to assign MLT dimensions to electric charge and taking charge as an additional basic quantity. But here, again, the distinction between charge and quantity is not recognized, leading to incorrect dimensions for electric current. These dimensions are stated as Q/T , the space-time equivalent of which is $1/s$, instead of the correct s/t . Both the MLT and the MLTQ systems of dimensional assignment are thus wrong in almost every electric and magnetic application, and they serve no useful purpose.

In our study of electrical fundamentals we were able to establish the correct dimensions of the electric quantities by using the mechanical dimensions as a base and taking advantage of the equivalence of mechanical and current phenomena. This approach is not feasible

in application to magnetism, but we have a good alternative, as our theory indicates that there is a specific dimensional relation between the magnetic quantities and the corresponding electric quantities, the dimensions of which we have already established.

The basic difference between electricity and magnetism is that electricity is one-dimensional whereas magnetism is two-dimensional. However, the various permutations and combinations of units of motion that account for the differences between one physical quantity and another are phenomena of only one scalar dimension, the dimension that is represented in the reference system. No more than this one dimension can be resolved into components by introduction of dimensions of space (or time). It follows that addition of a second dimension of motion to an electrical quantity takes the form of a simple unit of inverse speed, v . The dimensions of the magnetic quantity corresponding to any given electric quantity are therefore $\frac{1}{v}$ times the dimensions of the electric quantity. The dimensions thus derived for the principal magnetic quantities are shown in Table 30.

TABLE 30: *ELECTRICAL ANALOGS OF PRINCIPAL MAGNETIC QUANTITIES*

Electric		Magnetic	
t	dipole moment	t^2/s	dipole moment
$\frac{1}{s}$	charge	t^2/s^2	flux
$\frac{1}{s^2}$	potential	t^2/s^3	vector potential
$\frac{1}{s^3}$	flux density	t^2/s^4	flux density
$\frac{1}{s^3}$	field intensity	t^2/s^4	field intensity
t^2/s^2	resistivity	t^3/s^3	inductance
t^2/s^3	resistance	t^3/s^4	permeability

Here, then, we have a solid foundation for a critical analysis of magnetic relations, one that is free from the dimensional inconsistencies that have plagued magnetism ever since systematic investigation of magnetic phenomena was begun. In the next chapter we will apply the new understanding of magnetic fundamentals to an examination of magnetic quantities and units.

CHAPTER 20

Magnetic Quantities and Units

One of the major issues in the study of magnetism is the question as to the units in which magnetic quantities should be expressed, and the relations between them. “Since the first attempts to put its study on a quantitative basis,” says J. C. Anderson, “magnetism has been bedeviled by difficulties with units.”⁸³ As theories and mathematical methods of dealing with magnetic phenomena have come and gone, there has been a corresponding fluctuation in opinion as to how to define the various magnetic quantities, and what units should be used. Malcolm McCaig comments that, “with the possible exception of the 1940s, when the war gave us a respite, no decade has passed recently without some major change being made in the internationally agreed definitions of magnetic units.” He predicts a continuation of these modifications. “My reason for expecting further changes,” he says, “is because there are certain obvious practical inconveniences and philosophical contradictions in the SI system as it now stands.”⁸⁴

Actually, this difficulty with units is just another aspect of the dimensional confusion that exists in both electricity and magnetism. Now that we have established the general nature of magnetism and magnetic forces, our next objective will be to straighten out the dimensional relations, and to identify a consistent set of units. The ability to reduce all physical quantities to space-time terms has given us the tool by which this task can be accomplished. As we have seen in the preceding pages, this identification of the space-time relations plays a major part in the clarification of the physical situation. It enables us to recognize the equivalence of apparently distinct phenomena, to detect errors and omissions in statements of physical relationships, and to fit each individual relation into the total physical picture.

Furthermore, the verification process operates in both directions. The fact that all physical phenomena and relations can be expressed

in terms of space and time not only enables identifying the correct relations, but is also an impressive confirmation of the validity of the basic postulate which asserts that the physical universe is composed, in its entirety, of units of motion, an entity defined as a reciprocal relation between space and time.

The conventional treatment of magnetic phenomena employs the units of the mechanical and electrical systems so far as they are appropriate, and also, in some specialized applications, utilizes the same quantities under different names. For example, *inductance*, symbol L , is the term applied to the quantity involved in the production of an electromotive force in a conductor by means of variations in the current. The mathematical expression is

$$F = -L \frac{dI}{dt}$$

In space-time terms, the inductance is then

$$L = \frac{t}{s^2} \times \frac{t}{s} \times t = \frac{t^3}{s^3}$$

These are the dimensions of mass. Inductance is therefore equivalent to inertia. Because of the dimensional confusion in the magnetic area the inductance has often been regarded as being dimensionally equivalent to length, and the centimeter has been used as a unit, although the customary unit is now the henry, which has the correct dimensions. The true nature of the quantity known as inductance is illustrated by a comparison of the inductive force equation with the general force equation, $F = ma$.

$$F = ma = m \frac{dv}{dt} = \frac{m d^2s}{dt^2}$$

$$F = L \frac{dI}{dt} = L \frac{d^2q}{dt^2}$$

The equations are identical. As we have found, I (current) is a speed, and q (electric quantity) is space. It follows that m (mass) and L (inductance) are equivalent. The qualitative effects also lead to the same conclusion. Just as inertia resists any change in speed or velocity, inductance resists any change in the electric current.

Recognition of the equivalence of inductance and inertia clarifies some hitherto obscure aspects of the energy picture. An equivalent mass L moving with a speed I must have a kinetic energy $\frac{1}{2}LI^2$. We find experimentally that when a current I flowing through an inductance L is destroyed, an amount of energy $\frac{1}{2}LI^2$ does make its appearance. The explanation on the basis of conventional theory is that the energy is "stored in the electromagnetic field," but the identification of L with mass now shows that the expression $\frac{1}{2}LI^2$ is identical with the familiar expression $\frac{1}{2}mv^2$, and that, like its mechanical analog, it represents kinetic energy.

The inverse of inductance, t^3/s^3 , is *reluctance*, s^3/t^3 , the resistance of a magnetic circuit to the establishment of a magnetic flux by a magnetomotive force. As can be seen, this quantity has the dimensions of three-dimensional speed.

In addition to the quantities that can be expressed in terms of the units of the other classes of phenomena, there are also some magnetic quantities that are peculiar to magnetism, and therefore require different units. As brought out in the preceding chapter, these magnetic quantities and their units are analogous to the electric quantities and units defined in Chapter 16, differing from them only by reason of the two-dimensional nature of magnetism, which results in the introduction of an additional t/s term into each quantity.

The basic magnetic quantity, *magnetic charge*, is not recognized in current physical thought, but an equivalent quantity, *magnetic flux*, is used instead of charge, as well as in other applications where flux is the more appropriate term. The space-time dimensions of this quantity are the dimensions of electric charge, t/s , multiplied by the factor t/s that relates magnetism to electricity: $\text{t}/\text{s} \times \text{t}/\text{s} = \text{t}^2/\text{s}^2$. In the cgs system, magnetic flux is expressed in maxwells, a unit equivalent to 10^{-8} volt-sec. The SI unit is the weber, equivalent to the volt-sec. The justification for deriving the basic magnetic unit from an electric unit, the volt, can be seen when this derivation is expressed in space-time terms: $\text{t}/\text{s}^2 \times \text{t} = \text{t}^2/\text{s}^2$. These are the correct dimensions for magnetic charge or flux.

The natural unit of magnetic flux is the product of the natural unit of electric potential, 9.31146×10^8 volts, and the natural unit of time,

1.520655×10^{-16} seconds, and amounts to 1.41595×10^{-7} volt-sec, or webers. The natural units of other magnetic quantities can similarly be derived by combination of previously evaluated natural units.

Magnetic flux density, symbol B , is magnetic flux per unit area. The space-time dimensions are $t^2/s^2 \times 1/s^2 = t^2/s^4$. The units are the gauss (cgs) or tesla (SI). *Magnetic potential* (also called vector potential), like electric potential, is charge divided by distance, and therefore has the space-time dimensions $t^2/s^2 \times 1/s = t^2/s^3$. The cgs unit is the maxwell per centimeter, or gilbert. The SI unit is the weber per meter.

Since conventional physical science has never established the nature of the relation between electric, magnetic, and mechanical quantities, and has not recognized that an electric potential is a force, the physical relations involving the potential have never been fully developed. Extension of this poorly understood potential concept to magnetic phenomena has then led to a very confused view of the relation of magnetic potential to force and to magnetic phenomena in general.

As indicated above, the vector potential is the quantity corresponding to electric potential. The investigators working in this area also recognize what they call a *magnetic scalar potential*, which they define as $B \, ds/\mu$, where B is the magnetic flux density, t^2/s^4 , s is space and μ is a quantity with the dimensions t^3/s^4 that will be defined shortly. The space-time dimensions of the scalar potential are thus $t^2/s^4 \times s \times s^4/t^3 = s/t$. The so-called scalar potential is therefore a speed, equivalent to an electric current, a conclusion that agrees with the units, amperes, in which this quantity is measured. W. J. Duffin comments that it is not easy to put a physical interpretation on magnetic scalar potential.⁸⁵ The space-time dimensions of this quantity explain why. A potential (that is, a force) equivalent to a speed is a physical contradiction. The scalar potential is merely a mathematical construction without physical significance.

As indicated earlier, the magnetic quantities thus far defined are derived from the quantities of the mechanical and electrical systems. The units derived from the electrical system are related to the corresponding units of that system by the dimensions t/s , because of the two-dimensional nature of magnetism. Most of the

other magnetic quantities in common use are similarly derived, and all quantities of this set are therefore dimensionally consistent with each other and with the mechanical and electrical quantities previously defined in this and the preceding volume. But there are some other magnetic quantities that have been derived empirically, and are not consistent with the principal set of magnetic quantities or with the defined quantities in other fields. It is the existence of inconsistencies of this kind that has led to the conclusion of some physicists, expressed in a statement quoted in Chapter 9, that a consistent system of dimensions of physical quantities is impossible.

Analysis of this problem indicates that the difficulty, as far as magnetism is concerned, is mainly due to incorrect treatment of the dimensions of *permeability*, symbol μ , a quantity that enters into these and other magnetic relationships. The permeability of the great majority of substances is unity, or a close approximation thereto. The numerical results of magnetic measurements on these substances therefore give no indication of its existence, and there has been a tendency to overlook it, except where some collateral relation makes it clear that there are missing dimensions. But its field of application is actually very wide, as our theoretical development indicates that permeability is the magnetic equivalent of electrical resistance. It has the space-time dimensions of resistance, t^2/s^3 , multiplied by the factor $1/s$ that relates magnetism to electricity, the result being t^3/s^4 .

One of the empirical results that has contributed to the dimensional confusion is the experimental finding that *magnetomotive force* (MMF), or *magnetomotence*, is related to the current (I) by the expression $MMF = nI$, where n is the number of turns in a coil. Since n is dimensionless, this empirical relation indicates that the dimensions of magnetomotive force are the same as those of electric current. The SI unit of MMF has therefore been taken as the ampere.

It was noted in Chapter 9 that the early investigators of electrical phenomena attached the name "electromotive force" to a quantity that they recognized as having the characteristics of a force, an identification that we now find to be correct, notwithstanding the denial of its validity by most present-day physicists. A somewhat similar situation exists in magnetism. The early investigators in this

area identified a magnetic quantity as having the characteristics of a force, and gave it the name “magnetomotive force”. The prevailing view that this quantity is dimensionally equivalent to electric current contradicts the conclusion of the pioneer investigators, but here, again, our finding is that the original conception of the nature of this quantity is correct, at least in a general sense. Magnetomotive force, we find, is the magnetic (two-dimensional) analog of the one-dimensional quantity known as force. It has the dimensions of force, t/s^2 , multiplied by the factor $1/s$ that relates electricity to magnetism.

Dimensional consistency in magnetomotive force and related quantities can be attained by introducing the permeability in those places where it is applicable. Recognition of the broad field of applicability of this quantity has been slow in developing. As noted earlier, in most substances the permeability has the same value as if no matter is present, the reference level of unity, generally called the “permeability of free space.” Because of the relatively small number of substances in which the permeability must be taken into account, the fact that the dimensions of this quantity enter into many magnetic relations was not apparent in most of the early magnetic experiments. However, a few empirical relations did indicate the existence of such a quantity. For example, one of the important relations discovered in the early days of the investigation of magnetism is Ampère’s Law, which relates the intensity of the magnetic field to the current. The higher permeability of ferromagnetic materials had to be recognized in the statement of this relation. Permeability was originally defined as a dimensionless constant, the ratio between the permeability of the ferromagnetic substance and that of “free space.” But in order to make the mathematical expression of Ampère’s Law dimensionally consistent, some additional dimensions had to be included. The texts that define permeability as a ratio assign these dimensions to the numerical constant, an expedient which, as pointed out earlier, is logically indefensible. The more recent trend is to assign the dimensions to the permeability, where they belong. In the cgs system these dimensions are abhenry/cm. The abhenry is a unit of inductance, t^3/s^3 , and the dimensions of permeability on this basis are $\text{t}^3/\text{s}^3 \times 1/s = \text{t}^3/\text{s}^4$, which agrees with the previous determination. The SI units henry/meter and newton/ampere² ($\text{t/s}^2 \times \text{t}^2/\text{s}^2 = \text{t}^3/\text{s}^4$) are likewise

dimensionally correct. The unit farad/meter has been used, but this unit is dimensionless, as capacitance, of which the farad is the unit, has the dimensions of space [in application to static electricity]. Using this unit is equivalent to the earlier practice of treating permeability as a dimensionless constant. McCaig is quite critical of the unit henry/meter. He makes this comment:

Most books now quote the units of μ_0 as henry per metre. Although this usage is now almost universal, it seems to me to be a howler...The henry is a unit of self or mutual inductance and it seems quite incongruous to me to associate a metre of free space with any number of henries. If one wishes to be silly, one can invent numerous absurdities of this kind, e.g., torque is measured in Nm or joule!⁸⁶

The truth is that these two examples of what McCaig calls dimensional “absurdities” are quite different. His objection to coupling inductance with length is a purely subjective reaction, an opinion that they are incompatible quantities. Reduction of both quantities to space-time terms shows that his opinion is wrong. As indicated above, the quotient henry/meter has the dimensions t^3/s^4 , with a definite physical meaning. On the other hand, if the dimensions of torque are so assigned that they are equivalent to the dimensions of energy, there is a physical contradiction, as a torque must operate through a distance to do work; that is, to expend energy. This situation will be given further consideration later in the present chapter.

Returning now to the question as to the validity of the empirical relation $MMF = nI$, it is evident from the foregoing that the error in this equation is the failure to include the permeability, which has unit value under the conditions of the experiments, and therefore does not appear in the numerical results. When the permeability is inserted, the equation becomes $MMF = \mu nI$, the space-time dimensions of which are $t^2/s^3 = t^3/s^4 \times s/t$. The dimensions t^2/s^3 , which are assigned to MMF on this basis, are the appropriate dimensions for the magnetic analog of electric force, as they are the dimensions of force, $1/s^3$, multiplied by $1/s$, the dimensional relation between electricity and magnetism. These are also the dimensions of magnetic potential, previously defined. The justification for considering these as separate quantities is questionable.

In our previous consideration of a magnetic quantity currently measured in amperes, the magnetic scalar potential, we found that the assigned dimensions are correct, but that the quantity has no physical significance. In the case of the magnetomotive force, also measured in amperes in current practice, the magnetic quantity called by this name actually does exist in a physical sense, and it is a kind of force, but the dimensions currently assigned to it are wrong.

As in the electric system, the *magnetic field intensity* is the potential gradient, and should therefore have the dimensions $t^2/s^3 \times 1/s = t^2/s^4$, the same dimensions that we found for the flux density. The cgs unit, the oersted, is one gilbert per centimeter, and therefore has the correct dimensions. However, the unit in the SI system is the ampere per meter, the space-time dimensions of which are $s/t \times 1/s = 1/t$. These dimensions have been derived from the ampere unit of MMF, and the error in the dimensions of that quantity is carried forward to the magnetic field intensity. Introducing the permeability corrects the dimensional error and makes the cgs and SI units dimensionally consistent.

Magnetic pole strength is a quantity defined as F/B , where F is the force that is exerted. Again the permeability dimensions should be included. The correct definition is $\mu F/B$, the space-time dimensions of which are $t^3/s^4 \times t/s^2 \times s^4/t^2 = t^2/s^2$. Pole strength is thus merely another name for magnetic charge, as we might expect.

The permeability issue also enters into the question as to the definition of *magnetic moment*. The quantity currently called by that name, or designated as the *electromagnetic moment* (symbol m), is defined by the experimentally established relation $m = nIA$, where n and I have the same significance as in the related expression for the magnetomotive force, and A is the area of the circuit formed by each turn of a coil. The space-time dimensions are $s/t \times s^2 = s^3/t$. The moment per unit volume, the *magnetization*, M , is $s^3/t \times 1/s^3 = 1/t$.

An alternate definition of the magnetic moment introduces the permeability. This quantity, which is called the *magnetic dipole moment* to distinguish it from the moment defined in the preceding paragraph, has the composition μnIA . The space-time dimensions are $t^3/s^4 \times s/t \times s^2 = t^2/s$. (The distinction is not always effective, as some authors—Duffin, for example—apply the dipole moment designation

to the s^3/t quantity.) The dipole moment per unit volume, called the *magnetic polarization*, has the dimensions t^2/s^4 . This quantity is therefore dimensionally equivalent to the flux density and the magnetic field intensity, and is expressed in the same units.

The question as to whether the permeability should be included in the “moment” affects other magnetic relations, particularly that between the flux density B and a quantity that has been given the symbol H . This is the quantity with the dimensions $1/t$ that, in the SI system, is called the field intensity, or field strength. Malcolm McCaig reports that “the name field for the vector H went out of fashion for a time,” and says that he was asked by publishers to use “magnetizing force” instead. But “the term magnetic field strength now seems to be in fashion again.”⁸⁷ The relation between B and H has supplied the fuel for some of the most active controversies in magnetic circles. McCaig discusses these controversial issues at length in an appendix to his book *Permanent Magnets in Theory and Practice*. He points out that there are two theoretical systems that handle this relationship somewhat differently. “Both systems have international approval,” he says, “but there are intolerant lobbies on both sides seeking to have the other system banned.” The two are distinguished by their respective definitions of the torque of a magnet. The Kennelly system uses the magnetic dipole moment (t^2/s), and expresses the torque as $T = mH$. The Sommerfeld system uses the electromagnetic moment (s^3/t) and expresses the torque as $T = mB$.

Torque is a product of force and distance, $t/s^2 \times s = t/s$. The space-time dimensions of the product mH are $t^2/s \times 1/t = t/s$. The equation $T = mH$ is thus dimensionally correct. The space-time dimensions of the product mB are $s^3/t \times t^2/s^4 = t/s$. So the equation $T = mB$ is likewise dimensionally correct. The only difference between the two is that in the Kennelly system the permeability is included in m , whereas in the Sommerfeld system it is included in B .

This situation emphasizes the importance of a knowledge of the space-time dimensions of physical quantities, particularly in determining the nature of the connection between one quantity and another. A mathematically correct statement of a physical relation is not necessarily a true statement, because at least some of the terms

of that relation must have physical dimensions (otherwise it would be merely a mathematical statement, not a physical statement), and if those dimensions are wrong, the statement itself is *physically* wrong, regardless of its mathematical accuracy. The dimensions constitute a description of the physical nature of the quantities to which they apply, and give the mathematical statement of each relation a physical meaning.

As matters now stand, this is not recognized by everyone. McCaig, for example, indicates, in his discussion, that he holds an alternate view, in which the dimensions are seen as merely a reflection of the method of measurement of the quantities. He cites the case of force, which, he says, could have been defined on the basis of the gravitational equation, rather than by Newton's second law, in which event the dimensions would be different.

The truth is that we do not have this option, because the dimensions are inherent in the physical relations. In any instance where two different derivations lead to different dimensions for a physical quantity, one of the derivations is necessarily wrong. The case cited by McCaig is a good example. The conventional dimensional interpretation of the gravitational equation is obviously incompatible with the accepted *definition* of force based on Newton's second law of motion. Force cannot be proportional to the second power of the mass, as required by the prevailing interpretation of the gravitational equation, and also proportional to the first power of the mass, as required by the second law. And it is evident that an interpretation of the gravitational equation that conflicts with the definition of force is wrong. Furthermore, this equation, as interpreted, is an orphan. The physicists have not been able to reconcile it with physical theory in general, and have simply swept the problem under the rug by assigning dimensions to the gravitational constant.

McCaig's comments about the dimensions of torque emphasize the need to bear in mind that a numerically consistent relation does not necessarily represent physical reality, even if it is also consistent dimensionally. Good mathematics is not necessarily good physics. An example is the definition of torque as Fs , the product of the force and the lever arm (a distance). The work of rotation is defined as the

product of the torque and the angle of displacement ϑ . The work is thus $Fs\vartheta$. But work is the product of a force and the distance through which the force acts. This distance, in rotation, is not ϑ , which is purely numerical, nor is it the lever arm, because the length of the lever arm is not the distance through which the force acts. The effective distance is $s\vartheta$. Thus the work is not $Fs \times \vartheta$ (torque $\times$ angle), but $F \times s\vartheta$ (force $\times$ distance). Torque is actually a force, and the lever arm belongs with the angular displacement, not with the force. Its numerical value has been moved to the force merely for convenience in calculation. Such transpositions do not affect the *mathematical* validity, but it should be understood that the modified relation does not represent physical reality, and *physical* conclusions drawn from it are not necessarily valid.

Reduction of the dimensions of all physical quantities to space-time terms, an operation that is feasible in a universe where all physical entities and phenomena are manifestations of motion, not only clarifies the points discussed in the preceding pages, but also accomplishes a similar clarification of the physical situation in general. One point of importance in the present connection is that when the dimensions of the various quantities are thus expressed, it becomes possible to take advantage of the general dimensional relation between electricity and magnetism as an aid in determining the status of magnetic quantities.

For instance, an examination in the light of this relation makes it evident that identification of the vector H as the magnetic field intensity is incorrect. The role of this quantity H in magnetic theory has been primarily that of a mathematical factor rather than an expression of an actual physical relation. As one textbook comments, "the physical significance of the vector H is obscure."⁸⁸ (This explains why there has been so much question as to what to call it.) Thus there has been no physical constraint on the assignment of dimensions to this quantity. The unit of H in the SI system is the ampere per meter, the dimensions of which are $\frac{1}{t} \times \frac{1}{s} = \frac{1}{t}$. It does not necessarily follow that there is any phenomenon in which H can be identified physically. In current flow, the quantity $\frac{1}{s}$ appears as power. Whether the quantity $\frac{1}{t}$ has a role of this kind in magnetism is not yet clear. In any event, H is not the magnetic field intensity,

and should be given another name. Some authors tacitly recognize this point by calling it simply the “H vector.”

As noted earlier, the magnetic field intensity has the dimensions t^2/s^4 , and is therefore equivalent to μH ($t^3/s^4 \times 1/t$) rather than to H . This relation is illustrated in the following comparison between electric and magnetic quantities:

FIELD INTENSITY OR FLUX DENSITY

Electric

$$E = V/s = t/s^2 \times 1/s = t/s^3$$

Potential per unit space

$$E = R/t = t^2/s^3 \times 1/t = t/s^3$$

Resistance per unit time

Magnetic

$$B = A/s = t^2/s^3 \times 1/s = t^2/s^4$$

Potential per unit space

$$\mu H = \mu/t = t^3/s^4 \times 1/t = t^2/s^4$$

Permeability per unit time

Ordinarily the electric field intensity is regarded as the potential per unit distance, the manner in which it normally enters into the static relations. As the tabulation indicates, it can alternatively be regarded as the resistance per unit time, the expression that is appropriate for application to electric current phenomena. Similarly, the corresponding magnetic quantity B or μH , can be regarded either as the magnetic potential per unit space or the permeability per unit time.

A dimensional issue is also involved in the relation between magnetization, symbol M , and magnetic polarization, symbol P . Both are defined as magnetic moment per unit volume. The magnetic moment entering into magnetization is s^3/t , and the dimensions of this quantity are therefore $s^3/t \times 1/s^3 = 1/t$, making magnetization dimensionally equivalent to H . The magnetic moment entering into the polarization is the one that is generally called the magnetic dipole moment, dimensions t^2/s . The polarization is then $t^2/s \times 1/s^3 = t^2/s^4$. Magnetic polarization is thus dimensionally equivalent to field intensity B . To summarize the foregoing, we may say that there are two sets of these magnetic quantities that represent essentially the same phenomena, and differ only in that one includes the permeability, t^3/s^4 , while the other does not. The following tabulation compares the two sets of quantities:

Magnetic moment	s^3/t	Dipole moment	$t^3/s^4 \times s^3/t = t^2/s$
Magnetization	$1/t$	Polarization	$t^3/s^4 \times 1/t = t^2/s^4$
Vector H	$1/t$	Field Intensity	$t^3/s^4 \times 1/t = t^2/s^4$

A point to be noted about these quantities is that the magnetic polarization is not the magnetic quantity corresponding to the electric polarization. The magnetic polarization is a magnetostatic quantity, with dimensions t^2/s^4 , and its electric analog would be an electrostatic quantity with dimensions t/s^3 . This is what electric polarization would be on the basis of the conventional theory of storage of electric charge in capacitors. But, as we saw in Chapter 15, the capacitor stores electric current, not electric charge. It has therefore been found necessary to introduce a term with the dimensions s^2/t into the mathematical relations, eliminating the electrostatic quantities; that is, reducing coulombs (t/s) to coulombs (s). The need for this mathematical adjustment is a verification of our conclusion that the electrical storage process does not involve any polarization in the electrostatic sense.

The magnetic quantities identified in the discussion in this chapter—the principal magnetic quantities, we may say—are listed in Table 31, with their space-time dimensions and their units in the SI system.

TABLE 31: *MAGNETIC QUANTITIES*

Quantity	SI Units	Dimensions
dipole moment	weber $\times$ meter	t^2/s
flux	weber	t^2/s^2
pole strength	weber	t^2/s^2
vector potential	weber/meter	t^2/s^3
MMF		t^2/s^3
flux density	tesla	t^2/s^4
field intensity		t^2/s^4
polarization	tesla	t^2/s^4
inductance	henry	t^3/s^3
permeability	henry/meter	t^3/s^4
magnetization	ampere/meter	$1/t$
vector H	ampere/meter	$1/t$
magnetic moment	ampere $\times$ meter ²	s^2/t
reluctance	1/henry	s^3/t^3

The magnetic scalar potential has been omitted from the tabulation, for the reasons previously given, together with a number of other quantities identified in the contemporary magnetic literature in connection with

individual magnetic phenomena that we are not examining in this volume, or in connection with special mathematical techniques utilized in dealing with magnetism. The dimensionally incorrect SI units for MMF and magnetic field intensity are likewise omitted.

There is a question as to how far we ought to go in attaching different names to quantities that have the same dimensions and are therefore essentially equivalent. It would appear that the primary criterion should be usefulness. It is undoubtedly useful to distinguish clearly between electric quantity (space) and extension space, but it is not so clear that this is true of the distinction between the various quantities with the dimensions t^2/s^4 , for example. The magnetic field intensity can be identified with these dimensions, by analogy with the electric field intensity. Perhaps there is some justification for distinguishing it from magnetic polarization, which has the same dimensions. Whether this is also true of the other t^2/s^4 quantities such as flux density and magnetic induction is somewhat questionable.

The mathematical treatment of magnetism has improved very substantially in recent years, and the number of dimensional inconsistencies of the kind discussed in the preceding pages is now relatively small compared to the situation that existed a few decades earlier. But the present-day theoretical treatment of magnetism tends to deal with mathematical abstractions, and to lose contact with physical reality. The conceptual understanding of magnetic phenomena therefore lags far behind the mathematical treatment. This is graphically illustrated in Table 32. The upper section of this tabulation shows the “corresponding quantities in electric and magnetic circuits,”⁸⁹ according to a current textbook, with the space-time dimensions of each quantity, as determined in the present investigation. The lower section shows the correct analogs (magnetic = electric $\times t/s$) in the three cases where a magnetic analog actually exists. Only two of the seven identifications in the textbook are correct, and in both of these cases the dimensions that are currently assigned to the magnetic quantity are wrong. As brought out in the preceding discussion, the permeability, which belongs in both the MMF and the magnetic field intensity, is omitted from these quantities in the SI system.

TABLE 32: CORRESPONDING QUANTITIES

Electric		Magnetic	
<i>From reference 89, with space-time dimensions added</i>			
s/t	Current	t^2/s^2	Magnetic flux
$1/st$	Current density	t^2/s^4	Magnetic induction
s^2/t^2	Conductivity	t^3/s^4	Permeability
t/s^2	Electromotance	t^2/s^3	Magnetomotance
t/s^3	Electric field intensity	t^2/s^4	Magnetic field intensity
s^3/t^2	Conductance	t^3/s^3	Permeance
t^2/s^3	Resistance	s^3/t^3	Reluctance
<i>Correct analogs (magnetic = electric $\times t/s$)</i>			
s/t	Current		no magnetic analog
$1/st$	Current density		no magnetic analog
s^2/t^2	Conductivity		no magnetic analog
t/s^2	Electromotance	t^2/s^3	Magnetomotance
t/s^3	Electric field intensity	t^2/s^4	Magnetic field intensity
s^3/t^2	Conductance		no magnetic analog
t^2/s^3	Resistance	t^3/s^4	Permeability

When the dimensions of the various magnetic quantities are assigned in accordance with the specifications in the preceding pages, these quantities are all consistent with each other, and with the previously defined quantities of the mechanical and electric systems. This eliminates the need for employing illegitimate artifices such as attaching dimensions to pure numbers. The numerical magnitudes of the existing valid magnetic relations have already been adjusted in previous practice to fit the observations, and are not altered by the dimensional clarification.

The dimensional clarification in the magnetic area completes the consolidation of the various systems of measurement into one comprehensive and consistent system in which all physical quantities and units can be expressed in terms that are reducible to space and time only. There are, of course, many specialized units that have not been considered in the pages of this and the preceding volume—such as the light year, a unit of distance; the electron-volt, a unit of energy; the atmosphere, a unit of pressure; and so on—but

the quantities measured in these units are the basic quantities, or combinations thereof, and their units are specifically related to the units of space and time, both conceptually and mathematically.

CHAPTER 21

Electromagnetism

The terms “electric” and “magnetic” were introduced in Volume I with the understanding that they were to be used as synonyms for “scalar one-dimensional” and “scalar two-dimensional” respectively, rather than being restricted to the relatively narrow significance that they have in common usage. These words have been used in the same senses in this volume, although the broad scope of their definitions is not as evident as in Volume I, because we are now dealing mainly with phenomena that are commonly called “electric” or “magnetic.” We have identified a one-dimensional movement of uncharged electrons as an *electric current*, a one-dimensional rotational vibration as an *electric charge*, and a two-dimensional rotational vibration as a *magnetic charge*. More specifically, the magnetic charge is a two-dimensional rotationally distributed scalar motion of a vibrational character. Now we are ready to examine some motions that are not charges, but have some of the primary characteristics of the magnetic charge; that is, they are two-dimensional directionally distributed scalar motions.

Let us consider a short section of a conductor, through which we will pass an electric current. The matter of which the conductor is composed is subject to gravitation, which is a three-dimensional distributed inward scalar motion. As we have seen, the current is a movement of space (electrons) through the matter of the conductor, equivalent to an outward scalar motion of the matter through space. Thus the one-dimensional motion of the current opposes the portion of the inward scalar motion of gravitation that is effective in the scalar dimension of the spatial reference system.

For purposes of this example, let us assume that the two opposing motions in this section of the conductor are equal in magnitude. The net motion in this scalar dimension is then zero. What remains of

the original three-dimensional gravitational motion is a rotationally distributed scalar motion in the other two scalar dimensions. Since this remaining motion is scalar and two-dimensional, it is *magnetic*, and is known as *electromagnetism*. In the usual case, the gravitational motion in the dimension of the current is only partially neutralized by the current flow, but this does not change the nature of the result; it merely reduces the magnitude of the magnetic effect.

From the foregoing explanation it can be seen that electromagnetism is the residue of the gravitational motion that remains after all or part of the motion in one of the three gravitational dimensions has been neutralized by the oppositely directed motion of the electric current. Thus it is a two-dimensional scalar motion perpendicular to the flow of current. Since it is the gravitational motion in the two dimensions that are not subject to the outward motion of the electric current, it has the inward scalar direction.

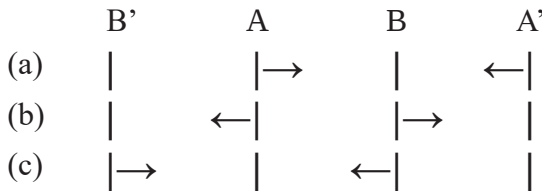
In all cases, the magnetic effect appears much greater than the gravitational effect that is eliminated, when viewed in the context of our gravitationally bound reference system. This does not mean that something has been created by the current. What has happened is that certain motions have been transformed into other types of motion that are more concentrated in the reference system, and energy has been brought in from the outside to meet the requirements of the new situation. As pointed out in Chapter 14, this difference that we observe between the magnitudes of motions with different numbers of effective dimensions is an artificial product of our position in the gravitationally bound system, a position that greatly exaggerates the size of the spatial unit. From the standpoint of the natural reference system, the system to which the universe actually conforms, the basic units are independent of dimensions; that is, $1^3 = 1^2 = 1$. But because of our asymmetric position in the universe, the natural unit of speed, $\frac{t}{s}$, takes the large value 3×10^{10} cm/sec, and this becomes a dimensional factor that enters into every relation between quantities of different dimensions.

For example, the c^2 term (the second power of 3×10^{10}) in Einstein's equation for the relation between mass and energy reflects the factor applicable to the two scalar dimensions that separate mass (t^3/s^3) from

energy ($\frac{1}{2}$). Similarly, the difference of one dimension between the two-dimensional magnetic effect and the three-dimensional gravitational effect makes the magnetic effect 3×10^{10} times as great (when expressed in cgs units). The magnetic effect is less than the one-dimensional electric effect by the same factor. It follows that the magnetic *unit* of charge, or emu, defined by the magnetic equivalent of Coulomb's law, is 3×10^{10} times as large as the electric unit, or esu. The electric unit, 4.80287×10^{-10} esu, is equivalent to 1.60206×10^{-20} emu.

The relative scalar directions of the forces between current elements are opposite to the directions of the forces produced by electric and magnetic charges, as shown in Fig. 23, which should be compared with Fig 22 of Chapter 19. The inward electromagnetic motions are directed toward the zero points from which the motions of the charges are directed outward. Two conductors carrying current in the same direction, AB or A'B, analogous to like charges, move toward each other, as shown in line (a) of the diagram, instead of repelling each other, as like charges do. Two conductors carrying current in the direction BA or B'A, as shown in line (c), also move toward each other. But conductors carrying current in opposite directions, AB' and BA', analogous to unlike charges, move away from each other, as indicated in line (b).

FIGURE 23



These differences in origin and in scalar direction between the two kinds of magnetism also manifest themselves in some other ways. In our examination of these matters we will find it convenient to consider the force relations from a different point of view. Thus far, our discussion of the rotationally distributed scalar motions—gravitational, electric, and magnetic—has been carried on in terms of the forces exerted by discrete objects, essentially point sources of the effects under consideration. Now, in electromagnetism,

we are dealing with continuous sources. These are actually continuous arrays of discrete sources, as all physical phenomena exist only in the form of discrete units. It would therefore be possible to treat electromagnetic effects in the same manner as the effects due to the more readily identifiable point sources, but this approach to the continuous sources is complicated and difficult. A very substantial simplification is accomplished by introduction of the concept of the *field* discussed in Chapter 12.

This field approach is also applicable to the simpler gravitational and electrical phenomena. Indeed, it is the currently fashionable way of handling all of these (apparent) interactions, even though the alternate approach is, in some ways, better adapted to the discrete sources. In examining the basic nature of fields we may therefore look at the gravitational situation, which is, in most respects, the simplest of these phenomena. As we saw in Chapter 12, a mass A has a motion AB toward any other mass B in its vicinity. This motion is inherently indistinguishable from a motion BA of atom B. To the extent that actual motion of mass A is prevented by its inertia or otherwise, the motion of object A therefore appears in the reference system as a motion of object B, constituting an addition to the actual motion of that object.

The magnitude of this gravitational motion of mass A that is attributed to mass B is determined by the product of the masses A and B, and by the separation between the two, as is the motion of mass B, if the scalar motion AB is regarded as a motion of both objects. It then follows that each spatial location in the vicinity of object A can be assigned a magnitude and a direction, indicating the manner in which a mass of unit size *would move* under the influence of the gravitational force of object A if it occupied that location. The assemblage of these locations and the corresponding force vectors constitute the gravitational field of object A. Similarly, the distribution of the motion of an electric or magnetic charge defines an electric or magnetic field in the space surrounding this charge.

The mathematical expression of this explanation of the field of a mass or charge is identical with that which appears in currently accepted physical theory, but its conceptual basis is entirely different.

The conventional view is that the field is “something physically real in the space”³² around the originating object, and that the force is physically transmitted from one object to the other by this “something.” However, as P. W. Bridgman concluded, after carrying out a critical analysis of this situation, there is no evidence at all to justify the assumption that this “something” actually exists.²⁹ Our finding is that the field is not “something physical.” It is merely a mathematical consequence of the inability of the conventional reference system to represent scalar motion in its true character. But this recognition of its true status as a mathematical expedient does not negate its usefulness. The field approach remains the simplest and most convenient way of dealing mathematically with magnetism.

The field of a magnetic charge is defined in terms of the force experienced by a test magnet. The field of a magnetic pole—one end of a long bar magnet, for example—is therefore radial. As can be seen from the description of the origin of electromagnetism in the foregoing paragraphs, the field of a wire carrying an electric current would also be radial (in two dimensions) if it were defined in terms of the force experienced by an element of the current in a parallel conductor. But it is customary to define the electromagnetic field on the magnetostatic basis; that is, by the force experienced by a magnet, or an electromagnet in the form of a coil, a solenoid, which produces a radial field similar to that of a bar magnet by means of its geometrical arrangement. When the field of a current-carrying wire is thus defined, it circles the wire rather than extending out radially. The force exerted on the test magnet is then perpendicular to the field, as well as to the direction of the current flow.

Here is a direct challenge to physical theory, an apparent violation of physical principles that apply elsewhere. It is a challenge that has never before been met. The physicists have not even been able to devise a plausible hypothesis. So they simply note the anomaly, the “strange” characteristics of the magnetic effect. “The magnetic force has a strange directional character,” says Richard Feynman, “at every instant the force is always at right angles to the velocity vector.”⁹⁰ It is likely, however, that this perpendicular relation between the direction of current movement and the direction of the force would not seem so strange if magnets interacted only with magnets and

currents with currents. In that event, the magnetic effect of current on current would still be “at right angles to the velocity vector,” but it would be in the direction of the field, rather than perpendicular to it, as the field would have to be defined in terms of the action of current on current. When there is interaction between current and magnet, the resultant force is perpendicular to the magnetic field; that is, to the field intensity vector. A test magnet in an electromagnetic field does not move in the direction of the field, as would be expected, but moves in a perpendicular direction.

Notice how strange the direction of this force is. It is not in line with the field, nor is it in the direction of the current. Instead, the force is perpendicular to both the current and the field lines.⁹¹

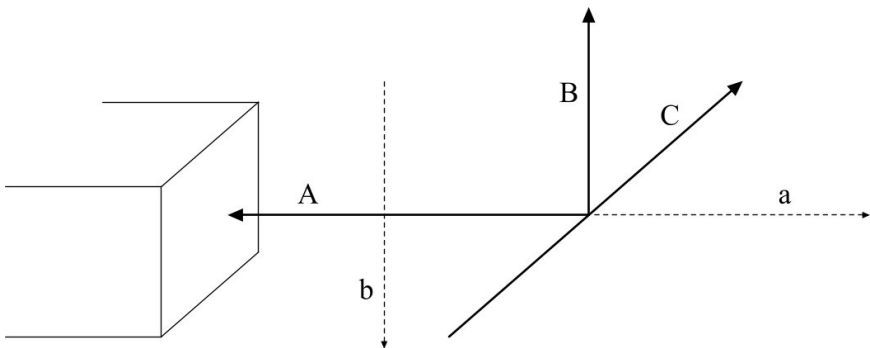
The use of the word “strange” in this statement is a tacit admission that the reason for the perpendicular direction is not understood in the context of present-day physical theory. Here, again, the development of the theory of the universe of motion provides the missing information. The key to an understanding of the situation is a recognition of the difference between the scalar direction of the motion (force) of the magnetic charge, which is outward, and that of the electromagnetic motion, which is inward.

The motion of the electric current obviously has to take place in one of the scalar dimensions other than that represented in the spatial reference system, as the direction of current flow does not normally coincide with the direction of motion of the conductor. The magnetic residue therefore consists of motion in the other unobservable dimension and in the dimension of the reference system. When the magnetic effect of one current interacts with that of another, the dimension of the motion of current A that is parallel with the dimension of the reference system coincides with the corresponding dimension of current B. As indicated in Chapter 13, the result is a single force, a mutual force of attraction or repulsion that decreases or increases the distance between A and B. But if the interaction is between current A and magnet C, the dimensions parallel to the reference system cannot coincide, as the motion (and the corresponding force) of the current A is in the inward scalar direction, while that of the magnet C is outward.

It may be asked why these inward and outward motions cannot be combined on a positive and negative basis with a net resultant equal to the difference. The reason is that the inward motion of the conductor A toward the magnet C is also a motion of C toward A, since scalar motion is a mutual process. The outward motion of the magnet is likewise both a motion of C away from A and a motion of A away from C. It follows that these are two separate motions of both objects, one inward and one outward, not a combination of an inward motion of one object and an outward motion of the other. It then follows that the two motions must take place in different scalar dimensions. The force exerted on a current element in a magnetic field, the force aspect of the motion in the dimension of the reference system, is therefore perpendicular to the field.

These relations are illustrated in Fig. 24. At the left of the diagram is one end of a bar magnet. This magnet generates a magnetostatic (MS) field, which exists in two scalar dimensions. One dimension of any scalar motion can be so oriented that it is coincident with the dimension of the reference system. We will call this observable dimension of the MS motion A, using the capital letter to show its observable status, and representing the MS field by a solid line. The unobservable dimension of motion is designated b, and represented by a dashed line.

FIGURE 24



We now introduce an electric current in a third scalar dimension. As indicated above, this is also oriented coincident with the dimension of the reference system, and is designated as C. The current generates an electromagnetic (EM) field in the dimensions a and b perpendicular

to C. Since the MS motion has the outward scalar direction, while the EM motion is directed inward, the scalar dimensions of these motions coincident with the dimension of the reference system cannot be the same. The dimensions of the EM motion are therefore B and a; that is, the observable result of the interaction between the two types of magnetic motion is in the dimension B, perpendicular to both the MS field A and the current C.

The comment about the “strange” direction of the magnetic force quoted above is followed by this statement: “Another strange feature of this force” is that “if the field lines and the wire are parallel, then the force on the wire is zero.” In this case, too, the answer to the problem is provided by a consideration of the distribution of the motions among the three scalar dimensions. When the dimension of the current is C, perpendicular to the dimension A of the motion represented by the MS field, the EM field is in scalar dimensions a and B. We saw earlier that the observable dimensions of the inward EM motion and the outward MS motion cannot be coincident. Thus the EM motion in dimension a is unobservable. It follows that the motion in scalar dimension B, the dimension at right angles to both the current and the field has to be the one in which the observable magnetic effect takes place, as shown in Fig. 24. However, if the direction of the current is parallel to that of the magnetic field, the scalar dimensions of these motions (both outward) are coincident, and only one of the three scalar dimensions is required for both motions. This leaves two unobservable scalar dimensions available for the EM motion, and eliminates the observable interaction between the EM and MS fields.

As the foregoing discussion brings out, there are major differences between magnetostatics and electromagnetism. Present-day investigators know that these differences exist, but they are unwilling to recognize their true significance because current scientific opinion is committed to a belief in the validity of Ampère’s nineteenth century hypothesis that all magnetism is electromagnetism. According to this hypothesis, there are small circulating electric currents—“Ampèrian currents”—in magnetic materials whose existence is assumed in order to account for the magnetic effects.

This is an example of a situation, very common in present-day science, in which the scientific community continues to accept, and

build upon, hypotheses which have been revised so drastically to accommodate new information that the essence of the original hypothesis has been totally negated. It should be realized that there is no empirical support for Ampère's hypothesis. The existence of the Ampèrian currents is simply assumed. But today no one seems to have a very clear idea as to just *what* is being assumed. Ampère's hypothetical currents were miniature reproductions of the currents with which he was familiar. However, when it was found that individual atoms and particles exhibit magnetic effects, the original hypothesis had to be modified, and the Ampèrian currents are now regarded as existing within these individual units. At one time it appeared that the assumed orbital motion of the hypothetical electrons in the atoms would meet the requirements, but it is now conceded that something more is necessary. The current tendency is to assume that the electrons and other sub-atomic particles have some kind of a spin that produces the same effects as translational motion. The following comment from a 1981 textbook shows how vague the "Ampèrian current" hypothesis has become.

At the present time we do not know what goes on inside these basic particles [electrons, etc.], but we expect their magnetic effects will be found to be the result of charge motion (spinning of the particle, or motion of the charges within it).⁹²

Ampère's hypothesis was originally attractive because it explained one phenomenon (magnetostatics) in terms of another (electromagnetism), thereby apparently accomplishing an important simplification of magnetic theory. But it is abundantly clear by this time that there are major differences between the two magnetic phenomena, and just as soon as that fact became evident, the case in favor of Ampère's hypothesis crumbled. There is no longer any justification for equating the two types of magnetism. The continued adherence to this hypothesis and use of Ampèrian currents in magnetic theory is an illustration of the fact that there is inertia in the realm of ideas, as well as in the physical world.

The lack of any theory—or even a model—that would explain *how* either a magnetostatic or electromagnetic effect is produced has left magnetism in a confused state where contradictions and inconsistencies are so plentiful that none of them is taken very seriously. A somewhat similar situation was encountered in our

examination of electrical phenomena, particularly in the case of those issues affected by the lack of distinction between electric charge and electric quantity, but a much larger number of errors and omissions have converged to produce a rather chaotic condition in the conceptual aspects of magnetic theory. It is, in a way, somewhat surprising that the investigators in this field have made so much progress in the face of these obstacles.

As noted earlier, many of the physical quantities involved in electromagnetism are the same as those that enter into magnetostatic phenomena. These are quantities applicable to two-dimensional scalar relations, irrespective of the particular nature of the phenomena in which they participate. The electromagnetic units applicable to these quantities are therefore the same ones defined for magnetostatic phenomena in Chapter 20. Some of the relations between these quantities are also those of two-dimensional motions in general, rather than being peculiar to either magnetostatics or electromagnetism. More commonly, however, the relations involved in electromagnetism are analogous to those encountered in current electricity, as electromagnetism is a phenomenon of current flow rather than of magnetic charges.

One example is the force between currents. There is no electromagnetic relation analogous to the Coulomb equation. The theorists commonly use "current elements" for purposes of analysis, but such units obviously cannot be isolated. A simple interaction between two units, analogous to the interaction between two charges, therefore does not exist. Instead, the simplest electromagnetic interaction, the one that is used in defining the unit of current, the ampere, is the interaction between the magnetic forces of parallel wires carrying currents. Making use of the field concept, the advantage of which is quite evident in dealing with currents, we first define the magnetic field of one current in terms of the flux density, B . This quantity B has been found to be equal to $\mu_0 I / (2\pi s)$. The space-time dimensions of this expression are $t^3/s^4 \times s/t \times 1/s = t^2/s^4$, the correct dimensions of flux density. The force exerted by this field on a length ℓ of the parallel current-carrying wire is then $BI\ell$, dimensions $t^2/s^4 \times s/t \times s = 1/s^2$.

The expressions representing the two steps of this evaluation of the force can be consolidated, with the result that the force on wire B due

to the current in wire A is $\mu_0 I_A I_B \ell / (2\pi s)$. If the currents are equal this becomes $\mu_0 I^2 \ell / (2\pi s)$. There is some resemblance between this and an expression of the Coulomb type, but it actually represents a different kind of a relation. It is a magnetic (that is, two-dimensional) relation analogous to the electric equation $V = IR$. In this electric relation, the force is equal to the resistance times the current. In the magnetic relation the force on a unit length is equal to the permeability (the magnetic equivalent of resistance) times the square of the current.

The energy relations in electromagnetism have given the theorists considerable difficulty. A central issue is the question as to what takes the place of the mass that has an essential role in the analogous mechanical relations. The perplexity with which present-day scientists view this situation is illustrated by a comment from a current physics textbook. The author points out that the energy of the magnetic field varies as the second power of the current, and that the similarity to the variation of kinetic energy with the second power of the velocity suggests that the field energy may be the kinetic energy of the current. "This 'kinetic energy' of a current's magnetic field," he says, "suggests that it has something like mass."⁹³

The trouble with this suggestion is that the investigators have not been able to identify any electric or magnetic property that is "something like mass." Indeed, the most striking characteristic of the electric current is its immaterial character. The answer to the problem is provided by our finding that the electric current is a movement of units of space through matter, and that the effective mass of that matter has the same role in current flow as in the motion of matter through space. In the current flow we are not dealing with "something like mass," we are dealing with mass.

As brought out in Chapter 9, electrical resistance, R , is mass per unit time, t^2/s^3 . The product of resistance and time, Rt , that enters into the energy relations of current flow is therefore mass under another name. Since current, I , is speed, the electric energy equation, $W = RtI^2$, is identical with the equation for kinetic energy, $W = \frac{1}{2}mv^2$. The magnetic analog of resistance is permeability, with dimensions t^3/s^4 . Because of the additional $\frac{1}{s}$ term that enters into this two-dimensional quantity, the permeability is the mass per unit space,

a conclusion that is supported by observation. As expressed by Norman Feather, the mass “involves the product of the permeability of the medium and a configurational factor having the dimensions of a length.”⁹⁴ In some applications, the function of this mass term, dimensions t^3/s^3 , is clear enough to have led to its recognition under the name of inductance.

The basic equations employed in dealing with inductance are identical with the equations dealing with the motion of matter (mass) through space. We have already seen (Chapter 20) that the inductive force equation, $F = L \, dI/dt$, is identical with the general force equation, $F = m \, dv/dt$, or $F = ma$. Similarly, magnetic flux, which is dimensionally equivalent to momentum, is the product of inductance and current, LI , just as momentum is the product of mass and velocity, mv . It is not always possible to relate the more complex electromagnetic formulas directly to corresponding mechanical phenomena in this manner, but they can all be reduced to space-time terms and verified dimensionally. The theory of the universe of motion thus provides the complete and consistent framework for electric and magnetic relationships that has heretofore been lacking.

The finding that the one-dimensional motion of the electric current acting in opposition to the three-dimensional gravitational motion leaves a two-dimensional residue naturally leads to the conclusion that a two-dimensional magnetic motion similarly applied in opposition to gravitation will leave a one-dimensional residue, an electric current, if a conductor is appropriately located relative to the magnetic motion. This is the observed phenomenon known as electromagnetic induction. While they share the same name, this induction process has no relation to the induction of electric charges. The induction of charges results from the equivalence of a scalar motion AB and a similar motion BA , which leads to the establishment of an equilibrium between the two motions. As indicated above, electromagnetic induction is a result of the partial neutralization of gravitational motion by oppositely directed scalar motion in two dimensions.

This induction process is another of the aspects of electricity and magnetism that is unexplained in conventional science. As one textbook puts it,

Faraday discovered that whenever the current in the primary circuit 1 is caused to change, there is a current *induced* in circuit 2 while that change is occurring. This remarkable result is not in general derivable from any of the previously discussed properties of electromagnetism.⁹⁵

Here, again, the advantage of having at our disposal a *general* physical theory, one that is applicable to all subdivisions of physical activity, is demonstrated. Once the nature of electromagnetism is understood, it is apparent from the theoretical relation between electricity and magnetism that the existence of electromagnetic induction necessarily follows.

Since there is no freely moving magnetic particle corresponding to the electron, there is no magnetic current, but magnetic motion can be produced in a number of ways, each of which is a method of inducing electric currents or voltage differences. For example, the magnetic motion may originate mechanically. If a wire that forms part of an electrical circuit is moved in a magnetic field in such a way that the magnetic flux through the wire changes (equivalent to a magnetic motion), an electric current is induced in the circuit. A similar effect is produced if the magnetic field is varied, as, for instance, if it is generated by means of an alternating current.

The force aspect of the one-dimensional (electric) residual motion left by the magnetic motion in the electromagnetic induction process can, of course, be represented as an electric field, but because of the manner in which it is produced, this field is not at all like the fields of electric charges. As Arthur Kip points out, there is an “extreme contrast” between these two kinds of electric fields. He explains,

An induced emf implies an electric field, since it produces a force on a static charge. But this electric field, produced by a changing magnetic flux, has some properties which are quite different from those of an electrostatic field produced by fixed charges... the special property of this new sort of electric field is that its curl, or its line integral around a closed path, is not zero. In general, the electric field at any point in space can be broken into two parts, the part we have called electrostatic, whose curl is zero, and for which electrostatic potential differences can be defined, and a part which has a nonzero curl, for which a potential function is not applicable in the usual way.⁹⁶

While the substantial differences between the two kinds of electric fields are recognized in current physical thought, as indicated by this quotation, the reason for the existence of these differences has remained unidentified. Our finding is that the obstacle in the way of locating the answer to this problem has been the assumption that both fields are due to electric charges—static charges in the one case, moving charges in the other. Actually, the differences between the two kinds of electric fields are easily accounted for when it is recognized that the processes by which these fields have been produced are entirely different. Only one involves electric charges.

The treatment of this situation by different authors varies widely. Some textbook authors ignore the discrepancies between accepted theory and the observations. Others mention certain points of conflict, but do not follow them up. However, one of those quoted earlier in this volume, Professor W. J. Duffin of the University of Hull, takes a more critical look at some of these conflicts, and arrives at a number of conclusions which, so far as they go, parallel the conclusions of this work quite closely, although, of course, he does not take the final step of recognizing that these conflicts invalidate the foundations of the conventional theory of the electric current.

Like Arthur Kip (reference 96), Duffin emphasizes that the electric field produced by electromagnetic induction is quite different from the electrostatic field. But he goes a step farther and recognizes that the agency responsible for the existence of the field, which he identifies as the electromotive force (emf), must also differ from the electrostatic force. He then raises the issue as to what contributes to this emf. "Electrostatic fields cannot do so,"¹³ he says. Thus the description that he gives of the electric current produced by electromagnetic induction is completely non-electrostatic. An emf of non-electrostatic origin causes a current I to flow through a resistance R . Electric charges play no part in this process. "No charge accumulates at any point," and "no potential difference can be meaningfully said to exist between any two points."⁹⁷

Duffin evidently accepts the prevailing view of the current as a movement of charged electrons, but, as indicated in a previously quoted statement (reference 13), he realizes that the non-electrostatic

force (emf) must act on the “carriers of the charges” rather than on the charges. This makes the charges superfluous. Thus the essence of his findings from observation is that the electric currents produced by electromagnetic induction are non-electrostatic phenomena in which electric charges play no part. These are the currents of our ordinary experience, those that flow through the wires of our vast electrical networks.

In the course of the discussion of electricity and magnetism in the preceding pages we have identified a number of conflicts between the results of observation and the conventional “moving charge” theory of the electric current, the theory presented in all of the textbooks, including Duffin’s. These conflicts are serious enough to show that the current cannot be a flow of electric charges. Now we see that the ordinary electric currents with which the theory of current electricity deals are definitely non-electrostatic; that is, electric charges play no part in them. The case against the conventional theory of the current is thus conclusive, even without the new information made available by the development reported in this work.

CHAPTER 22

Magnetic Materials

The discussion of static magnetism in Chapter 19 was addressed to the type of two-dimensional rotational vibration known as ferromagnetism. This is the magnetism known to the general public, the magnetism of permanent magnets. As noted in that earlier discussion, ferromagnetism is present in only a relatively small number of substances, and since this was the only type of magnetism known to the early investigators, magnetism was considered to be some special kind of a phenomenon of limited scope. This general belief undoubtedly had a significant influence on the thinking that led to the conclusion that magnetism is a by-product of electricity. More recently, however, it has been found that there is another type of magnetism that is much weaker, but is common to all kinds of matter.

For an understanding of the nature of this second type of static magnetism one needs to recall that the basic rotation of all material atoms is two-dimensional. It follows from the previously developed principles governing the combination of motions that a two-dimensional vibration (charge) can be applied to this two-dimensional rotation. However, unlike the ferromagnetic charge, which is independent of the motion of the main body of the atom, this charge on the basic rotation of the atom is subject to the electric rotation of the atom in the third scalar dimension. This does not alter the vibrational character of the charge, but it distributes the magnetic motion (and force) over three dimensions, and thus reduces its effective magnitude to the gravitational level. To distinguish this type of charge from the ferromagnetic charge we will call it an *internal* magnetic charge.

As we have seen, the numerical factor relating the magnitudes of quantities differing by one scalar dimension, in terms of cgs units, is 3×10^{10} . Thus the effective magnetic motion is only 3×10^{10} times

the magnitude of the gravitational motion. The corresponding factor applicable to the interaction between a ferromagnetic charge and an internal magnetic charge is the square root of 3×10^{10} , which amounts to 1.73×10^5 . The internal magnetic effects are thus weaker than those due to ferromagnetism by about 10^5 .

The scalar direction of the internal magnetic charge, like that of all other electric and magnetic charges thus far considered, is outward. All magnetic (two-dimensional) rotation of atoms is also positive (net displacement in time) in the material sector of the universe. But the motion in the third scalar dimension, the electric dimension, is positive in the Division I and II elements and negative in the Division III and IV elements. As explained in Chapter 19, the all-positive magnetic rotations of the material sector have a polarity of a different type that is related to the directional distribution of the magnetic rotation. If an atom of an electropositive element is viewed from a given point in space—from above, for example—it is observed to have a specific magnetic rotational direction, clockwise or counterclockwise. The actual correlation with north and south has not yet been established, but for present purposes we may call the end of the atom that corresponds to the clockwise rotation its north pole. This is a general relation applying to all electropositive atoms. Because of the reversals at the unit levels, the north pole of an electronegative atom corresponds to counterclockwise rotation; that is, this north pole occupies a position corresponding to that which is occupied by the south pole of an electropositive atom.

When electropositive elements are subjected to the field of a magnet, the orientation of the poles is the same in both the atoms and the magnet (which is similarly positive). The atoms of these elements therefore tend to orient themselves with their magnetic axes parallel to the magnetic field, and to move toward the stronger part of the field; that is, they are attracted by permanent magnets. Such substances are called *paramagnetic*. Electronegative elements, which have the reverse polarity, are oriented with the poles of their atoms opposite to those of a magnet. This puts like poles together, causing repulsion. These atoms therefore tend to orient themselves perpendicular to the magnetic field, and to move toward the weaker part of the field. Substances of this kind are called *diamagnetic*.

In present-day magnetic theory diamagnetism is regarded as a universal property of matter, the origin of which is unexplained. "All materials are diamagnetic,"⁹⁸ says one textbook. On this basis, paramagnetism or ferromagnetism, where they exist, simply overpower the basic diamagnetism. Our finding is that each substance is *either* paramagnetic or diamagnetic, depending on the scalar direction of the rotation in the electric dimension. Ferromagnetic substances are paramagnetic with an additional two-dimensional rotational vibration of the kind previously described.

All elements of the electropositive divisions I and II, except beryllium and boron, are paramagnetic. As in the case of other properties previously discussed, the positive preference carries over into some of the adjoining elements of Division III. All other elements of the electronegative divisions III and IV, except oxygen, are diamagnetic.

The abnormal behavior of some of the elements of Group 2A is a result of the small size of this 8-member group, which permits the constituent elements, in some instances, to function as members of the inverse division of the group. Boron, for example, is normally the third member of the positive division of Group 2A, but it can alternatively act as the fifth member of the negative division of this group. Boron and beryllium are the positive elements nearest to the negative division in this group, and therefore the most subject to whatever influences tend to cause the polarity reversal. Just why oxygen is the element of the negative division in which the polarity reversal takes place is not yet known.

As brought out in Volume I, all chemical compounds are combinations of electropositive and electronegative components. The presence of any significant amount of motion in time (space displacement) in a molecular structure prevents establishment of the positive magnetic orientation. All compounds, except those that are ferromagnetic, or heavily weighted with paramagnetic elements, are therefore diamagnetic. This overwhelming preference for diamagnetism in the compounds is probably what led to the currently accepted hypothesis of a universal diamagnetism.

The intensity of the magnetic effect in a magnetic material is measured in terms of magnetization, symbol M , which was defined

in Chapter 20. The magnetization and the intensity of the applied field are additive. Both therefore have the dimensions of magnetic field intensity, t^2/s^4 , but for historical reasons the field intensity is customarily identified with the vector H , which has the dimensions $1/\text{t}$. Since the magnetization must have the same dimensions as the field intensity, it is also expressed in terms of a unit with the $1/\text{t}$ dimensions. As we saw in Chapter 20, the actual physical quantities are μM and μH , rather than M and H , but the permeability, μ , entering into these definitions is the "permeability of free space," μ_0 , which has unit magnitude. The dimensional error therefore does not affect the numerical results of calculations.

From the foregoing, the net total magnetic field intensity, B , is the sum of $\mu_0 M$ and $\mu_0 H$. For some purposes it is convenient to express this quantity in terms of H only. This is accomplished by introducing the *magnetic susceptibility*, χ , defined by the relation $\chi = M/H$. On this basis, $B = (1 + \chi)\mu_0 H$.

As indicated earlier, the internal magnetic effects are relatively weak. The susceptibilities of both paramagnetic and diamagnetic materials are therefore low. Those of the diamagnetic substances are also independent of temperature. Some studies of the factors that determine the magnitude of the internal magnetic susceptibility were undertaken in the early stages of the theoretical investigation whose results are here being reported, and calculations of the diamagnetic susceptibilities of a number of simple organic compounds were included in the first edition of this work. These results have not yet been reviewed in the light of the more complete understanding of the nature of magnetic phenomena that has been gained in the past several decades, but there are no obvious inconsistencies, and some consideration of these findings will be appropriate at this time.

As would be expected, since the internal magnetic charge is a modification of the magnetic component of the rotational motion of an atom, the magnetic susceptibility is the reciprocal of the effective magnetic rotational displacement. There are, of course, two possible values of this displacement for most elements, but the applicable value is often indicated by the environment; that is, association with elements of low displacement generally means that the lower value

will prevail, and vice versa. Carbon, for instance, takes its secondary displacement, one, in association with hydrogen, but changes to the primary displacement, two, in association with elements of the higher groups.

Another source of variability is introduced by the fact that the susceptibility, like most other physical properties, has an initial level, and this level is also influenced by environmental factors. At the present stage of the investigation we are not able to evaluate these factors from purely theoretical premises, but they vary in a fairly regular way in the various families of compounds. We can therefore establish what we may call semi-theoretical values of the diamagnetic susceptibility of many relatively simple organic compounds with the aid of series relationships.

The experimental values of the susceptibility of these compounds vary over a substantial range. It was found, however, in the original investigation, that, except for certain differences in the initial levels, the diamagnetic susceptibility has the same value as a constant, which we are calling the *refraction constant*, that determines the index of refraction. The properties of radiation will not be covered in this volume, but the measurements of the refractive index are much more accurate than those of the magnetic susceptibility. It will therefore be desirable to use the refraction constant as a base in the calculation of the susceptibilities, and some explanation of the manner in which that constant was derived will be required.

Like the internal susceptibility, the refraction constant is the reciprocal of the effective magnetic rotational displacement, the total displacement minus the initial level. As in the case of the susceptibility, the determination of this constant is complicated by a variability in the initial levels, especially those of the most common elements in the organic compounds, carbon and hydrogen. For convenience, both in calculation and in emphasizing the series relationships, a value of the refraction constant is first calculated on the basis of what we may regard as "normal" values. The deviation of the constant from the normal value is then determined for each compound.

Table 33 shows the derivation of the refraction factors in three representative organic families of compounds.

TABLE 33: *INDEX OF REFRACTION (N-1)/D*

	Dev.	k_r	697 k_r		Observed
ACIDS	(O-.389	CO-.389	C-.778	H-.778)	
acetic acid	0	.511	.356	.354	.356
propionic	0	.564	.393	.391	.393
butyric	0	.600	.418	.415	.417
valeric	0	.625	.436	.434	
caproic	0	.644	.449	.448	
enanthic	0	.659	.459	.458	
caprylic	3	.675	.470	.472	
pelargonic	5	.687	.479	.478	
capric	7	.697	.486	.485	
hendecanoic	9	.705	.491	.491	
lauric	11	.713	.496	.500	
myristic	15	.724	.505	.502	
palmitic	19	.733	.511	.511	
stearic	23	.741	.516	.514	
PARAFFINS (C-.778 H-.889)					
propane	5	.834	.581	.582	
butane	3	.820	.572		
pentane	3	.818	.570	.570	
hexane	3	.816	.568	.568	.569
heptane	3	.814	.567	.567	.568
octane	3	.813	.567	.5655	
nonane	3	.812	.566	.565	
decane	3	.812	.566	.5645	
hendecane	3	.811	.565	.566	
dodecane	0	.807	.563	.563	
tridecane	0	.807	.562	.575	
tetradecane	0	.807	.562		
pentadecane	0	.807	.562	.5605	
hexadecane	0	.807	.562	.561	
heptadecane	0	.807	.562	.562	
octadecane	0	.807	.562	.562	

2-Me propane	5	.827	.576	.577	
2-Me butane	5	.823	.573	.573	
2-Me pentane	3	.816	.568	.566	
2-Me hexane	3	.814	.567	.567	
2-Me heptane	3	.813	.567	.5655	
ESTERS (O-.389 CO-.389 C-.778 H-.778)					
methyl formate	0	.511	.356	.353	
ethyl	0	.564	.393	.390	.392
propyl	0	.600	.418	.417	.419
butyl	0	.625	.436	.437	
amyl	0	.644	.449	.447	.452
hexyl	3	.664	.462	.463	
octyl	5	.687	.479	.479	
isopropyl	0	.600	.418	.419	
isobutyl	0	.625	.436	.437	.438
isoamyl	0	.644	.449	.449	
methyl acetate	-3	.556	.387	.385	.389
ethyl	-3	.593	.413	.413	.417
propyl	0	.625	.436	.433	.434
butyl	0	.644	.449	.447	.448
amyl	0	.659	.459	.456	.461
hexyl	3	.675	.470	.470	
heptyl	3	.685	.477	.478	
isopropyl	0	.625	.436	.433	
isobutyl	0	.644	.449	.447	.448
isoamyl	0	.659	.459	.458	.459
methyl propionate	-3	.593	.413	.412	
ethyl	-3	.619	.431	.430	.432
propyl	0	.644	.449	.447	
butyl	0	.659	.459	.458	
methyl butyrate	-3	.619	.431	.431	
ethyl	0	.644	.449	.447	
propyl	0	.659	.459	.458	
butyl	0	.671	.467	.463	.467
amyl	3	.685	.477	.477	

In the acids, for example, the normal rotational displacement of the oxygen atoms and the carbon atom in the CO group is 2, while that of the hydrogen atoms and the remaining carbon atoms is 1. The normal initial level is $\frac{2}{9}$ in all cases. The normal refraction factors of the individual rotational mass units are then 0.778 for the displacement 1 atoms, and half this value, or 0.389 for those of displacement 2. All of the acids from acetic (C_2) to enanthic (C_7) inclusive have normal initial levels (no deviations), and the differences in the individual refraction factors are due entirely to a higher proportion of the 0.778 units as the size of the molecule increases. The normal initial level in the corresponding hydrocarbons, however, is only $\frac{1}{9}$, and when the molecular chain becomes long enough to free some of the hydrocarbon groups at the positive end of the molecule from the influence of the acid radical at the negative end, these groups revert to their normal initial levels as hydrocarbons, beginning with the CH_3 end group and moving inward. In caprylic acid (C_8), the three hydrogen atoms in the end group have made the change, those in the adjoining CH_2 group do likewise in pelargonic acid (C_9), and as the length of the molecule increases still further the hydrogen in additional CH_2 groups follows suit. The deviations from the normal values (expressed in numbers of $\frac{1}{9}$ units per molecule) are shown in the first column of Table 33. The second column shows the refraction constants, k_r , calculated by applying the deviations in column 1 to the normal values. In columns 3 and 4 the product $0.697 k_r$ is compared with the quantity $(n-1)/d$, where n is the refractive index at the sodium D wavelength and d is the density. The refraction constant is related to the natural unit wavelength rather than to the wavelength at which the measurements were made, but the difference is incorporated in the factor 0.697 that is applied before the comparison with the values derived from observation. An explanation of the derivation of this factor and the reason for making the correlation in this particular manner would require more discussion of radiation than is appropriate in this volume, but the status of the calculated refraction constants as specific functions of the composition of the compounds is evident.

In the paraffins the initial levels increase with increasing length of the molecule rather than decreasing as in the acids. As brought out in Volume I, the hydrocarbon molecules are not the symmetrical

structures that their formula molecules would seem to represent. For example, the formula for propane, as usually expressed, is $\text{CH}_3\text{CH}_2\text{CH}_3$, which indicates that the two end groups of the molecule are alike. But the analysis of this structure revealed that it is actually $\text{CH}_3\cdot\text{CH}_2\cdot\text{CH}_2\cdot\text{H}$, with a positive CH_3 group at one end and a negative hydrogen atom at the other. This negative hydrogen atom has a zero initial level, and it exerts enough influence to eliminate the initial level in the hydrogen atoms of the two CH_2 groups, giving the molecule a total of 5 units of deviation from the normal initial level. When another CH_2 group is added to form butane, the relative effect of the negative hydrogen atom is reduced, and the zero initial level is confined to the CH_2H combination, with 3 hydrogen atoms. The deviation continues on this basis up to hendecane (C_{11}), beyond which it is eliminated entirely, and the molecule as a whole takes the normal 0.889 refraction constant.

Also shown in Table 33 is a representative sample of the monobasic esters, which, as would be expected of acid derivatives, follow the same pattern as the acids. The only new feature is the appearance of a -3 deviation in some of the lower compounds. This appears to be due to a reversal of the influences that are responsible for the additional positive deviations in the lower paraffins, an interpretation that is supported by the fact that both end groups of the esters are positive.

The objective of Table 33 is merely to show how the refraction constants that are used in the susceptibility calculations are derived from the molecular composition and structure, and the number of compounds listed has been limited to those required for this purpose. The refraction constants used in application to the greater number and variety of compounds included in Table 34, which shows the kind of results that are obtained from the susceptibility calculations, are determined in the same manner.

As noted earlier, the diamagnetic susceptibility of an organic compound is equal to its refraction constant with an adjustment for a difference in the initial levels. The magnetic initial level is generally the same as that in refraction except in certain groups in which the level is modified by some factor not yet specifically identified, but apparently geometric. In the compounds listed in Table 34, the CH_3 , CH_2OH , and OH end groups have initial levels $\frac{1}{9}$ unit higher, per

TABLE 34: *DIAMAGNETIC SUSCEPTIBILITIES*

	kr	DI	DI/m	Calc.	Observed		
PARAFFINS							
propane	.834	2.00	.077	.911	.919*	.898	
pentane	.818	2.00	.048	.866	.874*	.874	
hexane	.816	2.00	.040	.856	.865*	.858	.888
heptane	.814	2.00	.034	.848	.851*	.850	
octane	.813	2.00	.030	.843	.846*	.845	.872
nonane	.812	2.00	.027	.839	.843*	.843	
decane	.812	2.00	.024	.836	.842*	.839	
2-Me propane	.827	2.00	.059	.886	.890*	.888	
2-Me butane	.823	3.00	.071	.894	.893*	.892	
2-Me pentane	.816	3.00	.060	.875	.873*	.873	
2-Me hexane	.814	3.00	.052	.866	.861*	.860	.862
2-Me heptane	.813	3.00	.045	.857	.857		
2,2-di Me propane	.823	2.00	.048	.871	.875*	.874	
2,2-di Me butane	.816	3.00	.060	.876	.885*	.883	.885
2,2-di Me pentane	.814	3.00	.052	.866	.868*	.866	.869
2,3-di Me butane	.809	4.00	.080	.889	.885*	.883	.885
2,3-di Me pentane	.809	4.00	.069	.878	.873*	.873	.875
2,3-di Me hexane	.808	4.00	.061	.869	.865*	.865	
2,2,3-tri Me butane	.809	4.00	.069	.878	.882*	.878	.884
2,2,3-tri Me pentane	.808	4.50	.068	.876	.874*	.872	.874
ACIDS							
acetic acid	.511	0.50	.016	.527	.525*	.520	.535
propionic	.564	1.00	.025	.589	.586*	.578	.587
butyric	.600	1.00	.024	.624	.625*	.627	.636
valeric	.625	1.50	.025	.650	.655*		
caproic	.644	2.00	.031	.675	.676*	.676	
enanthic	.659	2.00	.027	.686	.680*	.680	
ALCOHOLS							
methyl alcohol	.599	1.00	.056	.655	.660	.650	.674
ethyl	.658	2.00	.077	.735	.728*	.717	.744
propyl	.686	2.00	.059	.745	.752*	.740	.766

butyl	.708	2.00	.048	.756	.763*	.743	.758
amyl	.722	2.00	.040	.762	.766*	.766	
hexyl	.730	2.50	.043	.773	.774*	.775	.805
octyl	.744	2.50	.034	.778	.777*	.788	
dodecyl	.761	3.00	.028	.792	.792		
isopropyl	.686	2.50	.074	.760	.762*	.759	
isobutyl	.708	3.00	.071	.779	.779*	.772	.798
isoamyl	.722	3.00	.060	.782	.782*	.782	.799

MONOBASIC ESTERS

methyl formate	.511	0.50	.016	.527	.533*	.518	.533
ethyl	.564	1.00	.025	.589	.580*	.580	.588
propyl	.600	1.00	.021	.621	.625*	.623	
butyl	.625	1.00	.018	.643	.644*	.645	
methyl acetate	.556	1.00	.025	.581	.575*	.570	.590
ethyl	.593	1.00	.021	.614	.614*	.607	.627
propyl	.625	1.00	.018	.643	.645*	.645	.651
butyl	.644	1.50	.023	.667	.666*	.663	.667
methyl propionate	.593	1.50	.031	.624	.624*	.614	.628
ethyl	.619	1.50	.027	.646	.651*	.644	.651
propyl	.644	1.50	.023	.667	.671*	.666	
methyl butyrate	.619	1.50	.027	.646	.650*	.645	.650
ethyl	.644	1.50	.023	.667	.669*	.667	.669
propyl	.659	2.00	.028	.687	.687*	.687	
isobutyl formate	.625	1.50	.027	.652	.654*	.654	
isoamyl	.644	2.00	.031	.675	.674*	.675	
isopropyl acetate	.625	2.00	.035	.660	.656*	.656	
isobutyl	.644	2.00	.031	.675	.676*	.676	
isoamyl	.659	2.00	.027	.686	.687*	.687	.690

DIBASIC ESTERS

ethyl oxalate	.546	1.00	.013	.559	.560*	.552	.554
propyl	.585	1.00	.011	.596	.605*	.600	
methyl malonate	.514	1.00	.014	.528	.528*	.520	
ethyl	.564	1.00	.012	.576	.578*	.573	.578
methyl succinate	.537	1.50	.019	.556	.558*	.555	
ethyl	.578	2.00	.021	.599	.604*	.600	

AMINES

butylamine	.774	1.00	.024	.798	.805*	.806	
amyl	.779	1.00	.020	.799	.796*	.795	
heptyl	.786	1.00	.015	.801	.808*	.808	
diethyl	.774	1.00	.024	.798	.777*	.776	.835
dibutyl	.788	1.00	.014	.802	.802*	.802	

CYCLANES

cyclopentane	.784	2.23	.056	.840	.844*	.843	
cyclohexane	.787	0.89	.019	.806	.810*	.785	.810
Me cyclohexane	.790	0.89	.016	.806	.804	.792	
Et cyclohexane	.788	1.44	.023	.811	.812		

BENZENES

benzene	.778	-3.11	-.074	.704	.702*	.698	.732
toluene	.782	-3.50	-.063	.719	.718*	.712	.734
o-xylene	.786	-3.50	-.055	.731	.733*	.733	
m-xylene	.786	-4.28	-.067	.719	.721*	.720	.743
p-xylene	.786	-3.89	-.061	.725	.723*	.722	
ethylbenzene	.782	-3.50	-.055	.727	.727*	.738	

unit of rotational mass, than the refraction levels. Interior CH_2 groups are subject to a similar modification, half as large ($\frac{1}{18}$ unit) at certain points, as the molecular chains lengthen. The sum of the individual differences in initial level, ΔI , is $m'/9$, where m' is the number of rotational mass units in the modified end groups of the molecule, plus half of the number of units in the modified interior groups, with appropriate adjustments in special cases.

The average difference in initial level for a molecule of rotational mass m is then $m'/9m$. In Table 34 this value, shown as ΔI_m , is applied to the refractive constants of representative groups of simple organic compounds to arrive at the internal magnetic susceptibilities. The corresponding values from observation are listed in the last three columns of the table. Values marked with asterisks are taken from a recent compilation.⁹⁹ Where no measurement was available from this source, a representative value from the earlier reports is shown in the same column. The last two columns shown the range of results

reported from the earlier measurements. In the normal paraffins the association between the CH_2 group and the lone hydrogen atom at the negative end of the molecule is close enough to enable the $\text{CH}_2\text{-H}$ combination to act as the end group. This means that there are 18 rotational mass units in the end groups of each chain. The value of DI for these compounds is therefore $18/9 = 2$. Branching adds more ends to the molecule, and consequently increases DI. The 2-methyl paraffins add one CH_2 end group, raising DI to 3, the 2,3-dimethyl compounds add one more, bringing this quantity up to 4, and so on. A very close association, similar to that in the $\text{CH}_2\text{-H}$ combination, modifies this general pattern. In 2-methyl propane, for instance, the CHCH_3 combination acts as an interior group, and the value of DI for this compound is the same as that of the corresponding normal paraffin, butane. The $\text{C}(\text{CH}_3)_2$ combination likewise acts as an interior group in 2,2-dimethyl propane, and as a unit with only one end group in the higher 2,2-dimethyl paraffins.

Each of the interior CH_2 groups with the higher initial level adds nine rotational mass units rather than the 8 corresponding to the group formula. This seems to indicate that in these instances a $\text{CH}_2\text{-CH}_2$ combination is acting geometrically as if it were $\text{CH}_3\text{-CH}$. In the ring compounds the CH_2 and CH groups take the normal 8 and 7 unit values respectively.

The behavior of the substituted chain compounds is similar to that of the paraffins, but there is a greater range of variability because of the presence of components other than carbon and hydrogen. The alcohols, a typical family of this kind, have a CH_3 group at one end of the molecule and a CH_2OH group at the other. The value of ΔI for the longer chains is therefore $26/9 = 2.89$. In the lower alcohols, however, the CH_2 portion of the CH_2OH group reverts to the status of an interior group, and ΔI drops to 2.00. The methyl alcohol molecule goes a step farther and acts as if it has only one end. A similar pattern can be seen in other organic families, such as the esters. Since we have found that the effective units of some of these compounds in certain of the phenomena previously examined are double formula molecules, it appears likely that the magnetic behavior of methyl alcohol and other compounds with similar characteristics can be attributed to the size of the effective molecule.

No similar studies of paramagnetic materials have yet been made. Unlike diamagnetism, paramagnetism is temperature dependent. For an explanation of this dependence we need to recall that magnetism is a motion. One of the significant advantages of recognizing its status as a motion is that its effect on other motions can be evaluated in terms of a direct addition or subtraction, rather than having to be approached circuitously by means of some hypothetical mechanism. Diamagnetism, which is motion in time (negative) has no connection with the thermal motion, which is motion in space (positive). But paramagnetism is positive, and has an imputed direction opposite to that of the thermal motion. Thus an increase in temperature reduces the paramagnetic effect.

The internal magnetism, which has been the principal subject of discussion thus far in the present chapter, is of interest primarily because of the light that it sheds on the nature and properties of magnetism in general. From a practical standpoint, "magnetism" is synonymous with ferromagnetism. No systematic study of ferromagnetism in the context of the theory of the universe of motion has yet been undertaken. There are, however, a few points about the place of this phenomenon in the general physical picture that should be noted.

Ferromagnetism exists only below a temperature, the *Curie point*, which is specific for each substance. Inasmuch as this type of magnetism is restricted to positive elements and some of their compounds, ferromagnetic materials are also paramagnetic, and exhibit their paramagnetic properties above the Curie temperature. In this range, the susceptibility is linearly related to the temperature, but the relation is inverse; that is, the relation is between temperature and $1/\chi$.

In one respect there is a significant difference between the magnetic susceptibility and most of the physical properties discussed in the earlier pages. The specific heat of any given substance, for instance, decreases with decreasing temperature, and reaches zero at a particular temperature level. There is no negative specific heat. Consequently, the specific heat of the individual atom is zero at all temperatures below this level. But magnetic forces act upon magnetic substances at all temperatures below the critical temperature, as well as above it. What we have here is a difference in the significance of the zero point.

As explained in Volume I, the true datum of physical activity, the natural zero, is unit speed, the speed of light. *Natural* physical magnitudes extend from this natural zero to the natural unit of speed in space (our zero) in one direction, and to unit speed in time (inverse speed) in the other. These two speed ranges are identical, except for the inversion. Most of the physical magnitudes with which we deal are in the range from our zero to the speed of light, but there are some quantities that extend beyond the natural zero levels. This introduces some modifying factors into the physical relations, as the natural zero levels are limiting magnitudes of the kind discussed in Chapter 17; that is, points at which an inversion of most physical properties takes place.

For example, a property such as thermal radiation that increases with the temperature up to the unit temperature level (the natural zero) does not continue to increase as the temperature rises still farther. Instead, as we will see in Volume III, it undergoes a decrease symmetrical with the increase that takes place between zero and unit temperature. A somewhat similar reversal occurs in the case of those properties that extend into the region inside unit space, the time region, as we have called it, because all changes in this region take place in time, while the associated space remains constant at the unit level.

Ferromagnetism is a phenomenon of the time region, and its natural zero point (the Curie temperature) is therefore a boundary between two dissimilar regions, rather than a center of symmetry, like the speed of light, the natural zero of speed. Instead of following the kind of a linear relation that is characteristic of the properties of the regions outside unit space, the relation of ferromagnetism to temperature has a more complex form due to the substitution of the spatial equivalent of time for actual space in this region where no change in actual space takes place.

No detailed studies in this area have yet been undertaken, but it seems evident that in the more regular elements the magnetization is subject to the $(1-x^2)^{1/2}$ relation that applies to other time region properties examined earlier, and to a square root factor, which may also be inter-regional. It can therefore be expressed as $M=k(1-T^2)^{1/4}$. If the magnetization is stated as a fraction of the initial magnetization,

and the temperature is similarly stated as a fraction of the Curie temperature, the constant k is eliminated, and the values derived from the equation apply to all substances that follow the regular pattern. Within the limits of accuracy of the experimental data, the reduced magnetizations thus calculated are in agreement with the empirical values, as reported by D. H. Martin.¹⁰⁰

Because the internal magnetic charge is applied against the basic rotational motion of the atom, its force is symmetrically distributed in the same manner as the gravitational force. But, as we have seen, ferromagnetism is a motion of an individual, specifically located, component of the atom. The directional distribution of the ferromagnetic force in the reference system is therefore determined by the atomic orientation. If each atom acted independently, the orientation of the atoms of an aggregate would be random, but, in fact, each magnetically charged atom exerts a force on its magnetic neighbors, tending to line up these neighboring atoms with its own magnetic directions. This orienting effect encounters mechanical resistance, and is ordinarily limited in scope. For this reason, and because the relation of each magnetic aggregate to its magnetic environment changes from time to time, the magnetic orientation of an aggregate is not usually uniform. Instead, the aggregate is subdivided magnetically into a number of sections, generally called "domains."

Ordinarily, the domains are randomly oriented, and the effective magnetic force is reduced by the distribution over the different directions. Application of an external field forces a reorientation of the atoms to conform with the direction of the field, the extent of which depends on the strength of the field. This reorientation concentrates the magnetic effect in the direction of the field, and results in an increase in the effective magnetic force, reaching a maximum, the saturation level, when the reorientation is complete.

CHAPTER 23

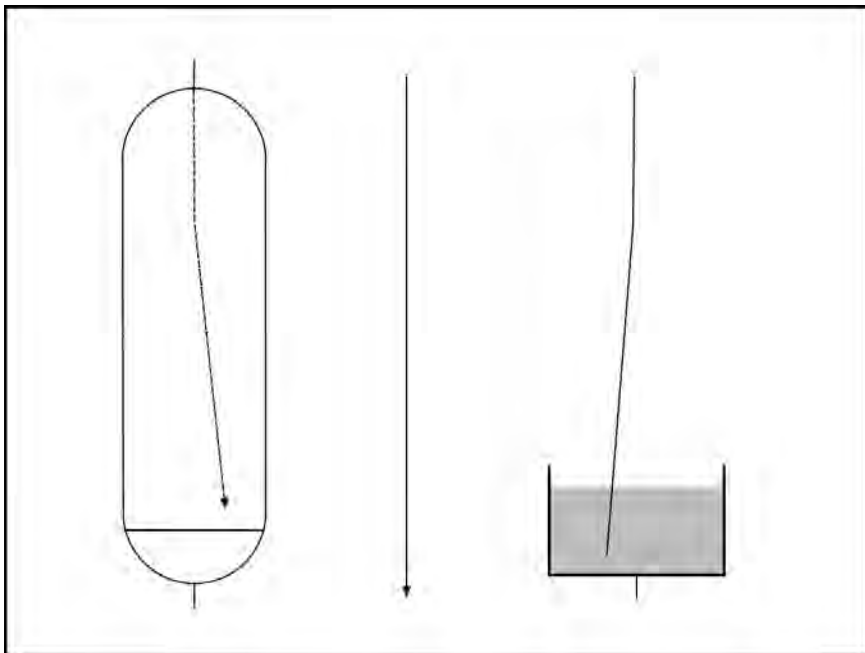
Charges in Motion

When a negative* charge is added to an electron, the net total scalar speed of the charged particle is zero. But since the electron rotation has the inward scalar direction, while the charge has the outward direction, the two motions take place in different scalar dimensions. Thus the electron does not act physically as a particle of zero speed displacement, but as an uncharged electron *and* a charge. A moving charged electron therefore has both magnetic properties (those of moving uncharged electrons) and electrostatic properties (those of charges).

The conventional view is that electrostatic phenomena are due to charges at rest, and magnetic phenomena are due to charges in motion. But, in fact, charges in motion have exactly the same electrostatic properties as charges at rest. “It is one of the remarkable properties of electric charge that it is invariant at all speeds,”¹⁰¹ says E. R. Dobbs. So the motion of the charges is not, in itself, sufficient to account for electromagnetism. Some additional process must come into operation in order to enable a charged particle to exhibit magnetic properties when in motion. Whether this additional process involves the charge or the particle—the “carrier of the charge,” as it was called in a statement previously quoted—is not specifically indicated observationally. Present-day theory simply assumes that all effects are due to the charges. But since there are “carriers,” these are obviously the moving entities. The charges have no motion of their own; they are carried. Even on the basis of conventional theory, therefore, the electromagnetic phenomena are due to the motion of the carriers, not motion of the charges. The development of electromagnetic theory in Chapter 21 now verifies this conclusion, and identifies the carriers of the charges as “bare” electrons.

As noted in Chapter 13, a flow of charged electrons through a conductor (a time structure) follows the same course as the flow of uncharged electrons. But the charged electrons have a property that their uncharged counterparts do not have. They can also move freely through the gravitational fields of extension space, producing electromagnetic phenomena that correspond to the effects of the flow of current in conductors. This is illustrated by an arrangement such as that shown in Fig. 25. In the center of the diagram is a wire through which a current is moving downward, as indicated by the arrow. (The conventional “direction of current flow” is opposite to the actual movement of the electrons, and is upward.) At the right is another conducting wire so arranged that a segment of the wire is hanging loose in a container filled with mercury. When a current is passed through this system in the same downward direction, the loose end of wire is attracted toward the center wire. At the left of the diagram is a vacuum tube through which a stream of electrons is also moving downward. This stream is attracted toward the center wire in the same way as the loose wire in the mercury container.

FIGURE 25



The movement of charged electrons through extension space is quite different in some other respects from the movement of uncharged electrons (space units) through matter. For instance, no electrical resistance is involved, and the motion therefore does not conform to Ohm's law. But the magnetic effect depends only on the neutralization of one dimension of a quantity of gravitational motion by the translational motion of the electrons, and from this standpoint the collateral properties of the motion are irrelevant. As long as the motion of the charged electrons takes place in a gravitational field, the requirement for the production of magnetic effects is met.

On the basis of the general principles applying to electromagnetic forces, as defined in Chapter 21, the magnetic force on a charged particle in a magnetic field is the product of the magnetic field intensity B and a motion combination with the dimensions s^2/l . The combination applicable to the motion of a charged particle, we find, is electric quantity q (measured as charge) multiplied by the particle velocity v . The force equation is then $F = Bqv$, with space-time dimensions $l/s^2 = l^2/s^4 \times s \times s/l$. The static force of the charge is $F = qE$, the dimensions of which are $l/s^2 = s \times l/s^3$.

The electrostatic forces between the charges (units of Q) are independent of the magnetic forces due to the movement of the electrons (units of q). The total force acting on a charged electron in a magnetic field is then $F = QE + Bqv$. Since Q and q are numerically equal, because each electron takes one unit of charge, this force expression can be written $F = q(E + Bv)$. The combined force is known as the Lorentz force. Lorrain and Corson comment on this force as follows:

The Lorentz force of Eq. 10-2 is intriguing. Why should $v \times B$ [velocity $\times$ magnetic field intensity] have the same effect as the electric field E ? Clearly from Eq. 10-2, the particle cannot tell whether it "sees" an E or a $v \times B$ term... Thus $v \times B$ is, somehow, an electric field intensity.¹⁰²

The authors then go on to say that the explanation is provided by the theory of relativity. But the space-time analysis shows that relativity has no bearing on this situation. From a physical standpoint, electric field intensity acts on a charged particle not as field intensity, but as a

quantity of $\frac{1}{s^3}$. Similarly, magnetic field intensity, $\frac{t^2}{s^4}$, acting against an electron moving with a velocity $\frac{s}{t}$ has the effect of a quantity of $(\frac{t^2}{s^4} \times \frac{s}{t})$; that is, a quantity of $\frac{1}{s^3}$. The magnitude of the physical result is the same in both cases.

This is not an unusual situation. On the contrary, it is common throughout all kinds of physical phenomena. The increase in temperature due to the addition of energy, for instance, depends entirely on the quantity of $\frac{1}{s}$ that is added to the thermal motion. It is immaterial whether that energy increment is in the form of kinetic energy, chemical energy, electrical energy, or any other form of $\frac{1}{s}$.

The effect of $\mathbf{v} \times \mathbf{B}$ does differ from that of $\mathbf{E}$ in direction, and the expression given for the Lorentz force is therefore valid only in vectorial form. The electric force $q\mathbf{E}$ acts in the direction of the field, and because the field is radial, the charges to which the force is applied "are accelerated, gaining kinetic energy."¹⁰³ The effect of the magnetic forces follows a different pattern. For the reasons explained in Chapter 21, the force exerted by a magnetic field on a moving electron is perpendicular to the field. As noted in the discussion of electromagnetism, this perpendicular direction of the force is an unexplained anomaly in present-day physical thought. "The strangest aspect of the magnetic force on a moving charge is the direction of the force,"¹⁰⁴ says a current textbook. When the origin of the magnetic field is understood, there is nothing strange about this direction. The scalar dimension of the motion of the electron is the dimension in which a portion of the gravitational motion is neutralized by the one-dimensional electron movement, and the residual two-dimensional motion necessarily exists in the two perpendicular dimensions.

The force aspect of this residual motion is also perpendicular to the magnetic field. If this is a magnetostatic field, it has the outward scalar direction, whereas the residual force has the inward scalar direction, and must therefore be in a different scalar dimension. If the field is electromagnetic, the forces are likewise in different dimensions, although the cause is different. As noted earlier, the motion of the uncharged electrons that constitute the electric current is in a scalar dimension other than that of the reference system. A freely moving

charged particle, on the other hand, is moving in the space, and therefore in the scalar dimension, of the reference system. The acceleration of an electron moving in a uniform magnetic field is thus perpendicular both to the field and to the direction of motion. Such an acceleration does not change the magnitude of the velocity; it merely changes the direction. Motion at constant speed with a constant acceleration at right angles to the velocity vector is motion in a circle. If the particle is also moving in a direction perpendicular to the plane of the circle, the path of motion is spiral.

Most of the empirical knowledge that has been gained with respect to the nature and properties of sub-atomic particles and cosmic atoms has been derived from observations of their motion in electric and magnetic fields. Unfortunately, the amount of information that can be obtained in this manner is very limited. A particularly significant point is that the experiments that can be made on electrons by the application of electric and magnetic forces are of no assistance to the physicists in their efforts to confirm one of their most cherished assumptions: the assumption that the electron is one of the basic constituents of matter. On the contrary, as pointed out in Chapter 18, the experimental evidence from this source shows that the assumed nuclear structure of the atom of matter which incorporates the electron is physically impossible.

The theory postulating orbital motion of negatively* charged electrons around a hypothetical positively* charged nucleus, developed by Rutherford and his associates after their celebrated experiments with alpha particles, collided immediately with one of the properties of the charged electrons. A charged object radiates if it is accelerated. Since the charge itself is an accelerated motion (for geometrical reasons), the force required to produce a given acceleration of the charge is less than that required to produce the same acceleration of the rotational unit. But the charge is physically associated with the rotational combination, and must maintain the same speed. The excess energy is therefore radiated away. This loss of energy from the hypothetical orbiting electrons would cause them to spiral in toward the hypothetical nucleus, and would make a stable atomic structure impossible.

This obstacle in the way of the nuclear hypothesis was never overcome. In order to establish the hypothetical structure as physically possible, it would be necessary (1) to determine just why an accelerated particle radiates, and (2) to explain why this process does not operate under the conditions specified in the hypothesis. Neither of these requirements has ever been met. Bohr simply assumed that the motion of the electrons is quantized and can take only certain specific values, thus setting the stage for all of the subsequent flights of fancy discussed in Chapter 18. The question as to whether the quantum assumption could be reconciled with the reasons for the emission of radiation by accelerated charges was simply ignored, as was the even more serious problem of accounting for the assumed coexistence of positive* and negative* charges at separations much less than those at which such charges are known to destroy each other. It should be no surprise that Heisenberg eventually had to conclude that the nuclear atom he helped to develop is not a physical particle at all, but is merely a “symbol,” that is, a mathematical convenience.

All of the foregoing discussion of the phenomena involving charges in motion has been carried out in terms of charged electrons. The same considerations apply, inversely in some respects, to charged positrons. Like the charged electrons, these positively* charged particles are capable of moving through space, and since their motion is outward, differing from that of the charged electrons only in rotational speed, they produce the same general kind of magnetic effect as the charged electrons. In the cosmic sector, the cosmic electric current is a flow of uncharged positrons through cosmic matter, and charged positrons moving through the cosmic gravitational fields in time have magnetic properties.

The rotational vibration that constitutes a charge may also be applied to other particles or to atoms. The charge on an atom or multi-unit particle and the unit of rotation that it modifies constitute a semi-independent component of that entity. The combination of charge and rotational unit remains as a constituent of the atom or particle, but vibrates independently, in the same manner as the magnetic motion combinations discussed in Chapter 19. Inasmuch as this vibrating combination has the same composition as a charged electron or

positron—a unit rotation modified by a unit rotational vibration—it has the same electric and magnetic properties.

The charges on atoms may be either positive* or negative*. As explained in Chapter 17, however, negative* ionization is confined to a relatively small number of elements because an atom must have a negative rotation in order to acquire a negative* (= positive) charge, and effective negative electric rotations are confined almost entirely to the elements of Division IV. On the other hand, any element can take a positive* charge. If the rotation in the electric dimension of the atom is negative, so that the positive* charge cannot be applied in this dimension, it can be applied to the rotation in one of the magnetic dimensions. The magnetic rotation is always positive in the material sector. It follows that while the mobile sub-atomic particles are predominantly negative*—that is, electrons—the freely moving (gaseous) ions are predominantly positive*.

The charged particles with which we have been concerned in the foregoing pages are electrically charged. Since there are also particles that are capable of taking magnetic charges, the question arises, Why do we not observe magnetically charged particles? The explanation can be found in the requirement that the net rotational displacement of a material atom or particle must be positive. The magnetic displacement, which is the larger component of the total, must therefore also be positive. This means that only negative magnetic charges can be applied to material particles.

The particles with magnetic rotational displacement are the neutron and the neutrino. The neutron has no electric displacement and only a single unit of magnetic displacement. Addition of an oppositely directed (negative) unit of charge therefore reduces the net displacement to zero, and terminates the existence of the particle. The neutrino has both electric and magnetic rotational components, and can therefore take a magnetic charge, but when it is in this charged condition it cannot move through space, for reasons that will be explained in Chapter 24, where the role of the charged neutrino in physical processes will be examined in detail.

The special contribution of magnetism to the verification of these significant consequences of the postulates that define the universe of

motion has been that, because of its intermediate position between the one-dimensional and three-dimensional phenomena, it, in a sense, ties the whole fabric of scalar motion theory together. Recognition of this point, early in the theoretical development, led to deferring consideration of magnetism until after the relations in the other major physical areas were quite firmly established. As a result, the investigation of magnetic phenomena is not as far advanced, particularly in quantitative terms, as the theoretical development in most of the other areas that have been covered.

There is also another factor that has limited the extent of coverage, one that is related to the objective of the presentation. This work is not intended as a comprehensive treatise on physics. It is simply an account of the results thus far obtained by development of the consequences of the postulates that define the universe of motion. In this development we are proceeding from the general principles expressed in the postulates toward their detailed applications. Meanwhile, the scientific community has been, and is, proceeding in the opposite direction, making observations and experiments, and working inductively from these factual premises toward increasingly general principles and relations. Thus the results of these two activities are moving toward each other. When the development of the Reciprocal System of theory reaches the point, in any field, where it meets the results that have been obtained inductively from observation and measurement, and there is substantial agreement, it is not necessary to proceed farther. Nothing would be gained by duplicating information that is already available in the scientific literature.

Obviously, the validity of existing theory in any particular area is one of the principal factors that determine just how far the new development needs to be carried in that area. As it happens, however, the previous work in magnetism, and to some extent in electricity as well, has followed along lines that are very different from those that are defined for us by the concept of a universe of motion, and the results of that previous work are, to a large extent, expressed in language that is altogether foreign to the manner in which our findings must necessarily be stated. This makes it difficult to determine just where we reach the point beyond which we are in agreement with previously existing theory. Whether the clarification

of the electric and magnetic relations in the special areas covered in the preceding pages will be sufficient, together with a translation of present-day theory into the appropriate language, to put electricity and magnetism on a sound theoretical footing, or whether some more radical reconstruction of theory will be required, is not definitely indicated as yet.

One step that clearly needs to be taken is to untangle electricity and magnetism from the theories and mathematics that apply specifically to radiation. An essential feature of this task will be to correct the many errors that have been introduced into electric and magnetic theory by means of the assumption that radiation is the *carrier* of electric and magnetic effects, an assumption whose detrimental effects on physical theory have been still further increased by the current tendency to extend it to gravitation. Radiation is outside the scope of this volume, and no extensive coverage of its properties will be undertaken in the present connection. However, the specific aspect of radiation that is relevant to the subject now under consideration is that, in the context of the stationary spatial reference system, a photon of radiation is a traveling quantity of energy. This was brought out in detail in Volume I in connection with the explanation of gravitation, in the course of which the characteristics of radiative energy were described in these words:

Radiation is a process whereby energy is transferred from a material aggregate at some location in space (or time) to other spatial (or temporal) locations. Each photon has a definite frequency of vibration and a correspondent energy content, hence these photons are essentially traveling units of energy. The emitting agency loses a specific amount of energy whenever a photon leaves. This energy travels through the intervening space (or time) until the photon encounters a unit of matter with which it is able to interact, whereupon the energy is transferred, wholly or in part, to this matter. At either end of the path the energy is recognizable as such, and is readily interchangeable with other types of energy.

Gravitational energy, it was pointed out, is not freely interchangeable with other forms of energy. At any specific location with respect to other masses, a mass unit possesses a definite amount of gravitational (potential) energy, and it is impossible to increase or

decrease this energy content at that location by conversion from or to other forms of energy. These facts are independent of the physical theory in whose context they are viewed. Even without the new information supplied by the Reciprocal System of theory it should be evident from the major differences between gravitation and radiation processes that the gravitational effects are not transmitted by any kind of radiation. The point of the present discussion is that electric and magnetic forces act in the same manner as gravitational forces, and the considerations that apply to the gravitational effects are likewise applicable to the electric and magnetic effects. There are not “carriers” of any of these effects.

It is true that radiation can be produced by electrical (as well as other) means, and this radiation can travel through space and produce physical effects at distant locations. But this process is totally unlike the way in which an electric or magnetic force acts, and the effects are entirely different. This radiation conforms to the description given in the paragraph from Volume I quoted above. The radiating object transfers energy to the radiation, and unless its energy content is continually replenished, it is eventually exhausted and the object ceases to radiate. An electric or magnetic charge, on the other hand, loses no energy in producing its characteristic effects.

The situation with respect to the object that is acted upon is not as clear-cut in the electric and magnetic phenomena as in gravitation because the properties of scalar motion that are involved in the process are unfamiliar, but obviously no energy can be transferred when none leaves the point of origin. This definitely established fact excludes any kind of an energy transfer process and provides a conclusive verification of the theoretical finding that the effects of electricity and magnetism on objects in their respective fields are separate and distinct from electromagnetic radiation and are phenomena of a very different character.

A striking result of the misunderstanding as to the role of radiation in physical processes has been that an almost mystical significance has been attributed to the quantity known as Planck’s constant (symbol h). This constant has become an indispensable tool for those who seek to evade the constraints that known physical laws and the rules of logic impose on theory construction. Niels Bohr,

for instance, called upon this constant in formulating a hypothesis that would enable the nuclear theory of the atom to circumvent the physical laws that stood in its way. In this case, the constant was used in defining a unit, or “quantum,” limiting the possible values of the speed of the hypothetical constituents of the atom. Later, when further difficulties emerged during the development of the details of the theory, Planck’s constant was again called upon to aid in defining additional “quantum numbers” based on “electron spin” and other hypothetical phenomena. “With the concept of spin the *universal constant of nature* h has almost become understandable,” says Walter R. Fuchs.¹⁰⁵

Werner Heisenberg employed the same mystical, “almost understandable” concept to add an air of mathematical legitimacy to his uncertainty principle, which he advanced as a justification for the inability of the quantum theories to derive specific numerical results in calculations involving events at the atomic level. According to his hypothesis, there is an unavoidable uncertainty of the order h/m in physical measurements. It follows that, in general, the quantum theories in whatever version happens to be in vogue at the moment, deal with probabilities rather than with specific values.

In view of the extraordinary importance ascribed to Planck’s constant in contemporary physical thought, it will probably come as a kind of a shock to find that in the universe of motion this constant has no physical significance at all; it is merely a mathematical relation whereby a quantity expressed in one manner can be expressed in a differing system of measurement. This constant is defined by the equation $W=h\nu$, where W is energy (work), ν is the radiation frequency and h is Planck’s constant, measured in erg-seconds. An erg, a unit of energy in the cgs system, is defined as $\text{gram} \times \text{cm}^2/\text{sec}^2$ which are indeed the dimensions of energy, since $t^3/s^3 \times s^2/t^2 = t/s$. An erg-second, energy $\times$ time, has the dimensions of t^2/s . Expressing frequency as $1/t$, the equation is dimensionally consistent, as can be seen by putting it into space-time terms: $t/s = t^2/s \times 1/t$.

However, the fact that a unit of radiation traveling through the space of the reference system exists in the time region, the region inside unit space, must also be taken into account. All physical

activity in this region is in time and the spatial equivalent of time (equivalent space, $\frac{1}{t}$) takes the place of space in physical relations. The space-time dimensions of the equation defining Planck's constant then become $t^2/(\frac{1}{t}) \times \frac{1}{t} = t^2$. Because of the location in the time region, the inter-regional ratio, 156.444 also applies. The value of the constant in sec^2/cm is then the second power of the natural unit of time divided by the interregional ratio. A further division by 2.236055×10^{-8} converts this value to erg-seconds:

$$h = \frac{(1.520655 \times 10^{-16})^2}{156.444 \times 2.236055 \times 10^{-8}} = 6.61028 \times 10^{-27} \text{ erg-sec}$$

The experimental value is reported as 6.62337×10^{-27} erg-sec. Here, again, there is a small discrepancy, two tenths of one percent, between the calculated and observed value. It is probably a result of the same small-scale factors that affect other calculations involving mass that were discussed in the preceding chapters.

Most of the other "universal constants" of physics have no more actual physical significance in the universe of motion than Planck's constant. The only numerical constants that enter into physical relations, when these relations are expressed in natural units, are what we may call structural constants, numbers that represent specific features of the physical situation. The number of effective dimensions is the structural constant most frequently encountered. The inter-regional ratio that entered into the evaluation of Planck's constant and had many other applications in the earlier pages, is a more complex magnitude of a similar nature likewise based on dimensional factors. In this connection it may be of interest to note that the fine structure constant, one of the physicists' "universal constants," is merely an alternate value of the inter-regional ratio that applies to certain types of interaction between matter and radiation. In these interactions the inter-regional ratio takes a value intermediate between the material value $(1 + \frac{2}{9})128$, and a radiation value.

Radiation is two-dimensional relative to the spatial reference system, as it is subject to the space-time progression in a dimension perpendicular to its inherent vibration. It therefore has the same $4 \times 4 \times 8$ distribution of time-region units as the two-dimensional basic rotations of the units of matter. But it does not have a sub-structure similar to

that responsible for the $\frac{2}{9}$ increment in the material inter-regional ratio, while it does have a negative initial level of one unit in the kind of interactions to which the fine structure constant applies. This initial level reduced the 16 two-dimensional units to 15, and the time-region total to 120. The composite inter-regional ratio applying to the interactions between radiation and matter is then $(120 \times 156.444)^{\frac{1}{2}} = 137.016$. A recent experimental measurement arrived at a value of 137.036.

This chapter concludes the discussion of magnetism as far as this subject will be covered in the present volume. Before turning to a different subject, it will be appropriate to make a few comments on the contents of the last five chapters and their relation to the physical situation in general.

Because the theory of the universe of motion, the detailed development of which is being described in these volumes, is new to the scientific community, and conflicts with many ideas and beliefs of long standing, the presentation in the several volumes of this series has a two-fold objective. It is designed not only to report the findings of the investigation based on the new theory, but also to provide the evidence that is required in order to confirm the validity of the findings. It therefore needs to be emphasized that the points brought out in the discussion of magnetism in these five chapters have made a very significant contribution to the mass of confirmatory evidence that is now available.

The particular importance of the magnetic evidence lies in the fact that the theory defines a specific dimensional relation between electricity and magnetism. It follows that whenever the theory identifies the nature of an electric phenomenon, this identification carries with it the assertion that there also exists a corresponding magnetic phenomenon, differing only in that it is two-dimensional, while the electric analog is one-dimensional.

Thus we find from the theory that there is a one-dimensional rotational vibration, identified as an electric charge, which has the space-time dimensions t/s and gives rise to a variety of electrostatic phenomena. According to the theory, it necessarily follows that there must be a two-dimensional rotational vibration, a magnetic charge, with the dimensions t^2/s^2 , that gives rise to an analogous variety of

magnetostatic phenomena. The observations verify the existence of a class of phenomena of this type, and an analysis of the dimensions of the magnetostatic quantities shows that they are, in fact, related to the corresponding electric quantities by the factor $\frac{1}{c}$, as required by the theory.

The dimensional interaction between electricity and magnetism is a particularly significant demonstration of the predictive power of the theory. We find from theory that gravitation is a three-dimensional scalar motion, and that an electric current is a one dimensional flow of units with the dimensions of space through the three-dimensional gravitating objects. From this it follows that the interaction should leave a two-dimensional scalar residue, oriented perpendicular to the current flow. Observations show that such a residue does exist, and that the process which leads to its existence can be identified with the phenomenon known as electromagnetism. It further follows from the same premises that the equivalent of a two-dimensional scalar motion through a three-dimensional gravitating object leaves a one-dimensional scalar motion as a residue. This interaction can be identified with the observed process known as electromagnetic induction, and the residue can be identified as an electric current.

The principal dimensional consequences that can be inferred from the theoretical identification of the electric current, electromagnetism, and gravitation with one, two, and three dimensions of scalar motion, respectively, are thus definitely correlated with observed electric and magnetic phenomena. But this is only the groundwork of a massive accumulation of evidence confirming the dimensional relations derived from theory.

Contemporary science places a great deal of emphasis on the predictive power of new theories. This is probably an overemphasis, as the ability of a theory to correlate existing information is as important as its ability to point the way to new information, and is becoming increasingly important as the "multitude of different parts and pieces" that now constitutes physical theory continues to expand. In any event, it should be recognized that deductions from the premises of a theory that identify hitherto unknown relations among known phenomena are predictions in the same sense as asserting the existence of a hitherto unknown phenomenon.

For example, the postulate that motion is the sole constituent of the physical universe carries with it the consequence that all physical quantities can be expressed in terms of space and time only. This is a prediction. The assertions as to the relation between electric and magnetic quantities discussed in the foregoing paragraphs are likewise predictions based on the same premises. The fact that the development of the consequences of the postulates of the theory of the universe of motion in the pages of this and the preceding volume has led to a complete and consistent system of space-time dimensions applicable to mechanical, electric, and magnetic quantities is a verification of these predictions.

The verification of this prediction is all the more significant because the possibility of arriving at any consistent system of dimensions, even with the use of four or five basic quantities, is denied by the majority of physicists.

In the past the subject of dimensions has been quite controversial. For years unsuccessful attempts were made to find “ultimate rational quantities” in terms of which to express all dimensional formulas. It is now universally agreed that there is no one “absolute” set of dimensional formulas.¹⁶

A similar prediction concerning the numerical values of these physical quantities is also implicit in the postulates. Since it is postulated that motion exists only in discrete units, it follows that the other physical quantities, all of which are either motions, combinations of motions, or relations between motions, likewise exist only in discrete units related to the units of the basic motion. This means that when the physical relations are correctly stated, they contain no numerical values other than those specifically identifying numbers of units, such as the atomic number, for example. The so-called “fundamental constants of physics” and the multitude of “disposable constants” that appear in relations such as the equations of state, will all be eliminated.

This fact that the values of the “fundamental constants” have no physical meaning in the context of the theory of the universe of motion contrasts sharply with the place of these constants in current scientific thought, where they are regarded as the critical magnitudes

that determine the nature of the universe. Paul Davies expresses the prevailing view in this statement:

...The gross structure of many of the familiar systems observed in nature is determined by a relatively small number of universal constants. Had these constants taken different numerical values from those observed, then these systems would differ correspondingly in their structure. What is specially interesting is that, in many cases, only a modest alteration of values would result in a drastic restructuring of the system concerned.¹⁰⁶

As we have seen in the pages of this and the preceding volume, some of these constants, the speed of light, the electron charge, etc., are natural units—that is, their true magnitude is unity—and the others are combinations of those basic units. The values that they take in the conventional measurement systems are entirely due to the arbitrary magnitudes of the units in which the measurements are expressed. The only way in which the constants could take the “different numerical values” to which Davies refers is by a modification of the measurement system. Such a change would have no physical meaning. Thus the possibility that he suggests in the quoted statement, and explores at length in the pages of his book *The Accidental Universe*, is ruled out by the unitary character of the universe. No physical relation in that universe is “accidental.” The existence of each relation, and the relevant magnitudes, are necessary consequences of the basic factors that define the universe as a whole, and there is no latitude for individual modification, except to the extent that selection among possible outcomes of physical events may be determined by probability considerations.

The clarification of these numerical relations to put them in terms of natural units is a gigantic task, and it is still far from being complete, but enough progress has been made, particularly in the fundamental areas, to make it evident that there is no serious obstacle in the way of continued progress toward the ultimate goal.

CHAPTER 24

Isotopes

While the magnetic charges involved in the phenomena that we recognize as magnetic all have the outward scalar direction, this does not mean that inward magnetic charges are non-existent. It is a result of the fact that the magnetic (two-dimensional) rotational displacement of material atoms is always inward. The principles governing the addition of motions, as set forth in Volume I, require charges to oppose the basic motions of the atoms in order to form stable combinations. The only *stable* magnetic charge is therefore the outward charge. However, inward charges may also be produced under appropriate conditions, and may continue to exist if their subsequent separation from the rotational combinations is forcibly prevented.

The events that take place during the beginning of the process of aggregation in the material environment were described in Volume I. As brought out there, the decay of the cosmic rays entering this environment produces a large number of massless neutrons, $M^{1/2-1/2-0}$. These are subject to disintegration into positrons, M^{0-0-1} , and neutrinos, $M^{1/2-1/2-(1)}$. Obviously, the presence of any such large concentration of particles of a particular type can be expected to have some kind of a significant effect on the physical system. We have already examined a wide variety of phenomena resulting from the analogous excess of electrons in the material environment. The neutrino is more elusive, and there is very little direct experimental information available concerning this particle and its behavior. However, the development of the Reciprocal System of theory has given us a theoretical explanation of the role of the neutrinos in physical phenomena, and we are now able to trace the course of events even where there are no empirical observations or data available for our guidance.

We can logically conclude that in some environments the neutrinos continue to exist in the uncharged condition in which they are originally formed, just as we find that the electron normally has no charge in the terrestrial environment. In this uncharged condition, the neutrino has a net displacement of zero. Thus it is able to move freely in either space or time. Furthermore, it is not affected by gravitation or by electrical or magnetic forces, since it has neither mass nor charge. It therefore has no motion relative to the natural system of reference, which means that from the standpoint of a stationary system of reference the neutrinos produced at any given location move outward at unit speed in the same manner as radiation. Each material aggregate in the universe is therefore exposed to a continuing flux of neutrinos, which may be regarded as a special kind of radiation.

Although the neutrino as a whole is neutral, from the space-time standpoint, because the displacements of its separate motions add up to zero, it actually has effective displacements in both the electric and magnetic dimensions. It is therefore capable of taking either a magnetic or an electric charge. Probability considerations favor the primary two-dimensional motion, and the charge acquired by a neutrino is therefore magnetic. This charge opposes the magnetic rotation, and since the rotation is inward the charge is directed outward. Inasmuch as this unit outward charge neutralizes the inward magnetic rotation, the only effective (unbalanced) unit of displacement of the charged neutrino is that of the inward negative rotation in the electric dimension. This charged neutrino is thus, in effect, a rotating unit of space, similar in this respect to the uncharged electron, and, as matters now stand, indistinguishable from it.

As a unit of space, the charged neutrino is subject to the same limitations as the analogous uncharged electron. It can move freely through the time displacements of matter, but it is barred from passage through open space, since the relation of space to space is not motion. Any neutrino that acquires a charge while passing through matter is therefore trapped. Unlike the charged electron, it cannot escape from the material aggregate by acquiring a charge. It must lose its charge in order to reach the neutral condition in which it is capable of moving through space. This is difficult to accomplish, as the conditions within the aggregate are favorable to producing charges rather than destroying

them. At first the proportion of neutrinos captured in passing through a newly formed material aggregate is probably small, but as the number of charged particles within the aggregate builds up, increasing what we may call the *magnetic temperature*, the tendency toward capture becomes greater. Being rotational in nature, the magnetic motion is not radiated away in the manner of the translational thermal motion, and the increase of the neutrino population is therefore a cumulative process. There will inevitably be some differences in the rate of build-up of this population by reason of local conditions, but in general the older a material aggregate becomes, the higher its magnetic temperature rises.

The charged neutrino, as a unit of space, is an addition to the space represented by the reference system, extension space, as we have called it. Where charged neutrinos are present, some of the atoms of matter exist, for a time, in the space of the neutrinos rather than in units of extension space, or in the space of the uncharged electrons that, as we have seen previously, are also present. The charged neutrinos are rotating relative to the spatial reference system, and they are consequently rotating relative to the systems of motions that constitute the material atoms, systems that are defined relative to the reference system. The outward rotational vibration (charge) of the spatial unit, the neutrino, is therefore equivalent to, and interchangeable with, an inward rotational vibration (charge) of the time structure, the atom. When the neutrino and the atom subsequently separate, there is a finite probability that the charge will remain with the atom.

The inward scalar direction of this two-dimensional atomic charge is the same as that of the two-dimensional atomic rotation. This fact that the rotational vibration of the atom induced by a magnetically charged neutrino is compatible with the basic magnetic (two-dimensional) inward rotation of the atom has a profound effect on the participation of this motion in physical processes. The ordinary magnetic charge is a foreign element in the material system, an outward motion in a system of inward motions. Magnetism therefore plays a detached role of relatively small importance in the local environment. The neutrino-induced rotational vibration, or charge, on the other hand, adds to the net rotational displacement (the mass) of the atom, and aside from being more dependent on conditions in

the environment, is fully coordinate with the basic atomic rotation. Instead of being a distinct *added* motion, this induced charge *modifies the magnitude* of the previously existing atomic rotation.

The presence of a concentration of charged neutrinos tending to produce inward rotational vibration of the atoms of an aggregate explains why an atom as a whole does not take an ordinary magnetic charge, and why these ordinary magnetic charges are confined to asymmetric atoms that have motion components which can vibrate independently of the main body. The outward motion cannot be initiated against the forces tending to produce inward motion.

In view of the very significant difference in behavior between the inward charge induced by the neutrinos and the ordinary outward magnetic charge, we will not use the term “magnetic charge” in application to the rotational vibration of the type we are now considering. Instead, we will call this a *gravitational charge*. Since the motion that constitutes this charge is a form of rotation, and is compatible with the atomic rotation, it adds to the net rotational displacement of the atom. However, there is only one rotating system in the neutrino, whereas the atom is a double system. The mass corresponding to the unit of gravitational charge is thus only half that of the unit of rotation (the unit of atomic number). For convenience, the smaller unit has been taken as the unit of atomic weight, or atomic mass. The primary atomic mass of a gravitationally charged atom is therefore $2Z + G$, where Z is the atomic number, and G is the number of units of gravitational charge.

In addition to the difference in the size of the units, the gravitational charge (rotational vibration) also has a relation to the atomic structure in general that is somewhat different from that of the full rotations. We will therefore distinguish between the *rotational mass* of the basic atomic rotation and the mass due to the gravitational charge, which we will call *vibrational mass*. The relation between the gravitational charge and the atomic rotation will have further consideration from the standpoint of the atomic structure in Chapters 25 and 26, and from the mass standpoint in Chapter 27.

Inasmuch as the gravitational charge is variable, the masses of the atoms of an element take different values, extending through a range that depends on the maximum size of the vibrational

mass G under the prevailing conditions. The different states that each element can assume by reason of the variable gravitational charge are identified as *isotopes* of the elements, and the mass on the $2Z+G$ basis is identified as the *isotopic mass*. As the elements occur naturally on the earth, the various isotopes of each element are almost always in the same, or nearly the same, proportions. Each element therefore has an average isotopic mass which is recognized as the atomic weight of that element. From the points brought out in the foregoing discussion, it is evident that the atomic weight thus defined is a reflection of the local neutrino concentration, the magnetic temperature, as we have called it, and does not necessarily have the same value in a different environment.

For reasons that will be explained in Chapter 26, the transfer of magnetic ionization from neutrino to atom is irreversible under terrestrial conditions. However, there are processes, to be described later, that gradually transform the vibrational mass into rotational mass. At a low magnetic temperature (concentration of charged neutrinos) most of the single gravitational charges are removed from the system by these processes before a second charge can be added. As the magnetic temperature increases, the frequency of magnetic ionization of atoms likewise increases because of the larger number of contacts. As a result, double or multiple ionization occurs in some atoms. Each aggregate thus has a *magnetic ionization level* analogous to the level of electric ionization previously discussed.

The degree of magnetic ionization of the individual elements depends not only on the magnetic temperature but also on the relative ability of those elements to absorb the neutrinos. This is a property of the individual units of time displacement. The effective magnetic ionization, the number of gravitational charges that are added to the atomic motion, therefore depends on the atomic mass as well as on the magnetic temperature. From the nature of the addition process we can deduce that at the unit ionization level each net unit of rotational displacement (atomic number) should be capable of acquiring one unit of gravitational charge (half the size of the atomic mass unit). But the atom exists in the time region, whereas the neutrino is not subject to the factors that apply to motion inside unit space. The relation between the charge and the atomic rotation is therefore

between m_v , the vibrational mass, and m_r^2 , the second power of the rotational mass. Furthermore, the atomic rotation in the time region is subject to the inter-regional ratio, 156.444. Denoting the magnetic ionization level as I , we then have the equilibrium relation

$$m_v = \frac{I m_r^2}{156.444} \quad (24-1)$$

In this equation the rotational mass, m_r , is expressed in the double units (units of atomic number) and the vibrational mass, m_v , in the single units (units of atomic weight).

The value of m_v thus derived is the number of units of gravitational charge (mass) that will normally be acquired by an atom of rotational mass m_r if raised to the magnetic ionization level I . It is quite obvious from the available empirical information that the magnetic ionization level on the surface of the earth is close to unity. A calculation for the element lead on the unit ionization basis, to illustrate the application of the equation, results in $m_v = 43$. Adding the 164 atomic weight units of rotational mass corresponding to atomic number 82, we arrive at a theoretical atomic weight of 207. The experimental value is 207.2.

This close agreement is not quite as significant as it appears to be. Actually there are stable isotopes of lead with isotopic masses ranging from 204 to 208. The explanation is that the value obtained from equation 24-1 is not necessarily the mass corresponding to the atomic weight, nor the isotopic mass of the most stable isotope. It is the center of a zone of isotopic stability. Because of the individual characteristics of the elements, the actual median of the stable isotopes and the measured atomic weight may be offset to some extent from this theoretical center of stability, but the deviation is generally small. In more than 60 percent of the first 92 elements it is only one unit, or none at all. Furthermore, the agreement is improving as more accurate measurements become available from experimental sources. In the nearly thirty years since the publication of the first edition of this work, in which the comparative values were tabulated, the accepted atomic weight of six elements has been changed by a significant amount, and in all of these cases the change has been in the direction of closer agreement with the theoretical values.

Table 35: *ATOMIC MASS EQUILIBRIUM VALUES*

Z	m				Z	m			
	m _v	Calc.	Obs.	Diff.		m _v	Calc.	Obs.	Diff.
1	.01	2	1	-1	47	14.12	108	108	0
2	.03	4	4	0	48	14.73	111	112.5	+1.5
3	.06	6	7	+1	49	15.35	113	115	+2
4	.10	8	9	+1	50	15.98	116	119	+3
5	.16	10	11	+1	51	16.63	119	122	+3
6	.23	12	12	0	52	17.28	121	128	+7
7	.31	14	14	0	53	17.96	124	127	+3
8	.41	16	16	0	54	18.64	127	131	+4
9	.52	19	19	0	55	19.34	129	133	+4
10	.64	21	20	-1	56	20.05	132	137	+5
11	.77	23	23	0	57	20.77	135	139	+4
12	.92	25	24	-1	58	21.50	138	140	+2
13	1.08	27	27	0	59	22.25	140	141	+1
14	1.25	29	28	-1	60	23.01	143	144	+1
15	1.44	31	31	0	61	23.78	146	145	-1
16	1.64	34	32	-2	62	24.57	149	150	+1
17	1.85	36	35.5	-0.5	63	25.37	151	152	+1
18	2.07	38	40	+2	64	26.18	154	157	+3
19	2.31	40	39	-1	65	27.01	157	159	+2
20	2.56	43	40	-3	66	27.84	160	162.5	+2.5
21	2.82	45	45	0	67	28.69	163	165	+2
22	3.09	47	48	+1	68	29.56	166	167	+1
23	3.38	49	51	+2	69	30.43	168	169	+1
24	3.68	52	52	0	70	31.32	171	173	+2
25	4.00	54	55	+1	71	32.22	174	175	+1
26	4.32	56	56	0	72	33.14	177	178.5	+1.5
27	4.66	59	59	0	73	34.06	180	181	+1
28	5.01	61	59	-2	74	35.00	183	184	+1
29	5.38	63	63.5	+0.5	75	35.96	186	186	0
30	5.75	66	65	-1	76	36.92	189	190	+1
31	6.14	68	70	+2	77	37.90	192	192	0
32	6.55	71	73	+2	78	38.89	195	195	0

33	6.96	73	75	+2	79	39.89	198	197	-1
34	7.39	75	79	+4	80	40.91	201	200.5	-0.5
35	7.83	78	80	+2	81	41.94	204	204	0
36	8.28	80	84	+4	82	42.98	207	207	0
37	8.75	83	85.5	+2.5	83	44.03	210	209	-1
38	9.23	85	88	+3	84	45.10	213	209	-4
39	9.72	88	89	+1	85	46.18	216	210	-6
40	10.23	90	91	+1	86	47.28	219	222	+3
41	10.74	93	93	0	87	48.38	222	223	+1
42	11.28	95	95.5	+0.5	88	49.50	226	226	0
43	11.82	98	98	0	89	50.63	229	227	-2
44	12.37	100	101	+1	90	51.78	232	232	0
45	12.94	103	103	0	91	52.93	235	231	-4
46	13.53	106	106.5	+0.5	92	54.10	238	238	0

Table 35 is an updated version of the original tabulation. The first column of the table gives the atomic number, The second column shows the value of m_v calculated from equation 24-1. Column 3 is the theoretical equilibrium mass, $2Z+G$, taken to the nearest unit, since the gravitational charge does not exist in fractional units. Column 4 is the observed atomic weight, also expressed in terms of the nearest integer, except where the excess is almost exactly one half unit. Column 5 is the difference between the observed and calculated values. The trans-uranium elements are omitted, as these elements cannot have (terrestrial) atomic weights in the same sense in which that term is used in application to the stable elements.

The width of the zone of stability is quite variable, ranging from zero for technetium and promethium to a little over ten percent of the rotational mass. The reasons for the individual differences in this respect are not yet entirely clear. One of the interesting, and probably significant, points in this connection is that the odd-numbered elements generally have narrower stability limits than the even-numbered elements. This and other factors affecting atomic stability will be discussed in Chapter 26. Isotopes that are outside the zone of stability undergo spontaneous modifications that have the result of moving the atom into the stable zone. The nature of these processes will be examined in the next chapter.

In addition to the limitation on its width, the zone of isotopic stability also has an upper limit due to the restrictions on the total rotation of the atom. It was established in Volume I that the maximum effective magnetic rotational displacement is four units. The elements of rotational group 4B have magnetic rotational displacements 4-4. By adding rotation in the electric dimension it is possible to build the total rotation up to 4-4-31, or the equivalent 5-4-(1), corresponding to atomic number 117, without exceeding the overall displacement maximum. But the next step brings the electric rotation up to the equivalent of the next unit of magnetic rotation. The effective magnetic rotation (that is, the total less the initial unit) is then four units in each magnetic dimension. As explained earlier, a displacement of four full magnetic units is equivalent to no displacement at all. Arrival at this point therefore terminates the rotation. The speed displacement reverts to the translational status. Element 118 is thus unstable, and will disintegrate promptly if it is formed. All rotational combinations above element 118 (rotational mass 236) are similarly unstable, whereas all elements below 118 are stable at a zero ionization level.

At a finite ionization level the corresponding vibrational mass is added to the rotational mass, and the 236 limit is reached at a lower atomic number. As indicated by Table 35, the equilibrium mass of uranium, atomic number 92, is 238 at the unit ionization level. This exceeds the 236 limit. Uranium and all elements above it in the atomic series are therefore unstable in an environment that is subject to this degree of ionization. The converse is not necessarily true; that is, it does not necessarily follow that all isotopes below the 236 limit are stable if they are within the zone of stability defined by the ratio of vibrational to rotational mass. At the magnetic temperature corresponding to the unit ionization level most atoms of an aggregate have one gravitational charge. But some have none, whereas others may possess two charges. The existence of a doubly charged atom has no observable physical consequences, other than the added mass, unless the second charge puts the total mass over the 236 limit. In that event, the atom will eventually disintegrate.

All of the factors that determine the extent of instability in the elements just below uranium in the atomic series have not yet been

identified, but, as would be expected, there is a general decrease in the tendency toward instability as the atomic number decreases. The lowest element that could theoretically become unstable by reason of acquisition of two gravitational charges is gold, element 79, for which the total mass with two units of charge is 238. However, the probability of the second ionization drops rapidly as we move down the atomic series, and while the first few elements below uranium are very unstable, the instability is negligible below bismuth, element 83.

As the magnetic ionization level rises, the stability limit decreases still further in terms of atomic number. It should be noted, however, that the rate of decrease slows down quickly. The first stage of ionization reduces the stability limit from 118 to 92, a difference of 26 in atomic number. The second unit of ionization causes a decrease of 13 atomic number units, the third only 8, and so on.

CHAPTER 25

Radioactivity

The ejection of positive or negative displacement by an atom that becomes unstable for one of the reasons discussed in the preceding pages will be identified as *radioactivity*, or *radioactive decay*, and the adjective *radioactive* will be applied to any element or isotope that is in the unstable condition. As brought out in Chapter 24, there are two distinct kinds of instability. Those elements whose atomic mass exceeds 236, either in rotational mass alone, or in rotational mass plus the vibrational mass added by magnetic ionization, are beyond the overall stability limit, and must reduce their respective masses below 236. In a fixed environment this cannot ordinarily be accomplished by modification of the vibrational mass alone, since the normal ratio of vibrational to rotational mass is determined by the prevailing magnetic ionization level. The radioactivity resulting from this cause therefore involves the actual ejection of mass and the transformation of the element into an element of a lower atomic number. The most common process is the ejection of one or more helium atoms, or alpha particles, and is known as *alpha decay*.

The second type of instability is due to a ratio of vibrational to rotational mass which is outside the zone of stability. In this case ejection of mass is not necessary; the required adjustment of the ratio can be accomplished by a process that converts vibrational mass into rotational mass, or vice versa, and thereby transforms the unstable isotope into another isotope within or closer to the zone of stability. The most common process of this kind is the emission of a beta particle, an electron or positron, together with a neutrino, and the term *beta decay* is applied.

In this work the alpha and beta designations will be used in a more general sense. All processes that result from instability due to exceeding the 236 mass limit (that is, all processes that involve the ejection of

primary mass) will be classified as alpha radioactivity, and all processes that modify only the ratio of vibrational mass to rotational mass will be classed as beta radioactivity. If it is necessary to identify the individual processes, such terms as β^+ decay, etc., will be employed. The nature of the processes will be specified in terms of the beta particle, and coincident emission of the appropriate neutrino should be understood.

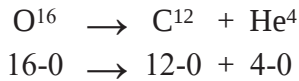
In analyzing these processes, which are few in number and relatively simple, the essential requirement is to distinguish clearly between the rotational and vibrational mass. For convenience we will adopt a notation in the form 6-1, where the first number represents the rotational mass and the second the vibrational mass. The example cited is the mass of the isotope Li^7 . A negative mass (space displacement) will be indicated by parentheses, as in the expression 4-(1), which is the mass of the isotope He^3 . This system is similar to the notation used for the rotational displacement in the different scalar dimensions, but there should be no confusion since one is a two-number system while the other employs three numbers.

Radioactive processes generally involve some adjustments of the secondary mass, but these are minor items that have not yet been studied in the context of the Reciprocal System of theory. They will not be considered in the present discussion, which will refer only to the primary mass, the principal component of the total.

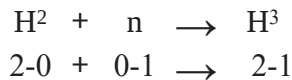
The composition of the motions of a stable isotope can be changed only by external means such as violent contact, absorption of a particle, or magnetic ionization, and the frequency of such changes is related to the nature of the environment, rather than to anything inherent in the structure of the isotope itself. An unstable isotope, on the other hand, is capable of moving toward stability on its own initiative by *ejecting* the appropriate motion or combination of motions. Consequently, each such process has a specific time pattern, subject to the probability relations.

The basic process of alpha radioactivity is the direct removal of rotational mass. Since each unit of rotational displacement is equal to two units of mass on the atomic weight scale, the effect of each step in this process is to decrease the rotational mass by $2n$ units. The rotational combination with $n = 1$ is the H^2 isotope, which is unstable because its total rotation is above the limit for either a single rotating

system, or an intermediate type of structure similar to that of the H^1 isotope, but is less than one double (atomic number) unit in each of the two rotating systems of the atomic structure. This H^2 isotope therefore tends either to lose displacement and revert to the H^1 status, or to add displacement and become a helium atom. The particle ejected in alpha radioactivity is the smallest stable double rotating system, in which $n=2$. Emission of this particle, the He^4 isotope, with mass components 4-0, results in a change such as



Since rotational vibration exists only as a modifier of rotation, there are no separate units of vibrational mass that can be added or subtracted directly in the manner of the alpha particle. But the mass of the compound neutron has the same single (atomic weight) unit value as the vibrational mass unit, and like the latter, it is a single rotating system (from the material standpoint). It is therefore interchangeable with the vibrational mass. In our numerical notation, it can be expressed as 0-1. This equivalence of the neutron mass and the unit of vibrational mass makes it possible to modify isotopes by adding or removing compound neutrons. Thus we may start with the mass two isotope of hydrogen, H^2 , and by adding a compound neutron obtain the mass three isotope, H^3 .



Beta radioactivity is a conversion process rather than an ordinary addition process. An isotope that is above the zone of stability has one or more units of magnetic displacement, $\frac{1}{2}-\frac{1}{2}-0$, in the form of rotational vibration, superimposed on units of the magnetic rotation of the atom. These vibrational units are only half the size of the rotational units. Addition of a second half-size unit to one of the combinations of unit rotation and unit rotational vibration is therefore required to produce an additional rotational unit. This cannot be accomplished by direct addition, as a rotational unit is not capable of accepting more than one vibrational unit. However, an unstable isotope is subject to influences that cause it to eject displacement.

(That is what makes it unstable.) An isotope above the stability zone ejects a cosmic neutrino, $(\frac{1}{2})-(\frac{1}{2})-1$ and an electron, $0-0-(1)$. This ejection is equivalent to addition of displacement $\frac{1}{2}-\frac{1}{2}-0$, the addition that is required to convert one of the half size vibrational units to a rotational unit. Because of the excess of material neutrinos in the material environment, the cosmic neutrino ejected from the atom is ultimately neutralized by one of the oppositely directed particles. Thus the net result of the ejection and its aftermath is the addition of a material neutrino to the atom and emission of an electron, which is equivalent to the addition of a positron. The same result can, of course, be accomplished by the direct addition of a material neutrino and a positron, but the scarcity of positrons in the local environment limits the availability of this direct process.

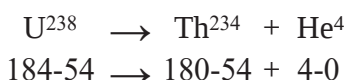
Neither of the ejected particles has any effective primary mass. No change in mass therefore takes place in this process (β^- radioactivity). The original isotope with rotational mass $2Z$ and vibrational mass n becomes an isotope with rotational mass $2(Z+1)$ —that is, an isotope of the next higher element—and vibrational mass $n-2$. The total mass of the combination of motions remains the same, but two units of vibrational mass have been converted to rotational mass, and the combination has moved closer to the zone of stability. If it is still outside that zone, the ejection process is repeated.

Where an isotope is below the zone of stability (deficient in vibrational mass) the process described in the foregoing paragraphs is reversed. In this process, β^+ radioactivity, a unit of rotational mass is converted to two units of vibrational mass by ejection of a material neutrino, $\frac{1}{2}-\frac{1}{2}-(1)$, together with a positron, $0-0-1$. The isotope of element Z , with rotational mass $2Z$ and vibrational mass n then becomes an isotope of element $Z-1$, with rotational mass $2(Z-1)$ and vibrational mass $n+2$.

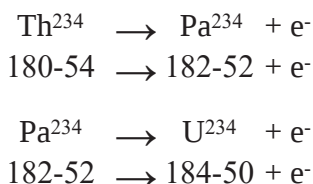
These are the basic radioactive processes. The actual course of events in any particular case depends on the initial situation. It may involve only one such event; it may consist of several successive events of the same kind, or a combination of the basic processes may be required to complete the transition to a stable condition. In natural beta radioactivity under terrestrial conditions a single beta emission is usually sufficient, as the unstable isotopes are seldom very far outside

the zone of beta stability. However, under some other environmental conditions the amount of radioactivity required in order to attain beta stability is very substantial, as we will see in Volume III.

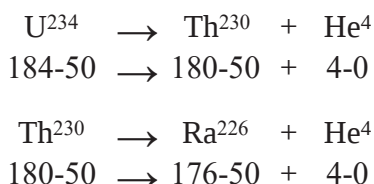
In natural alpha radioactivity, the mass that must be ejected may amount to the equivalent of several alpha particles even in the terrestrial environment. The loss of this rotational mass necessitates beta emission to restore the equilibrium between rotational and vibrational mass. Alpha radioactivity is thus usually a complex process. As an example, we may trace the steps involved in the radioactive decay of uranium. Beginning with U^{238} , which is just over the borderline of stability, and has the long half life of 4.5×10^9 years, the first event is an alpha emission.



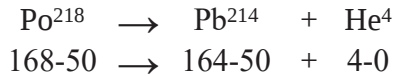
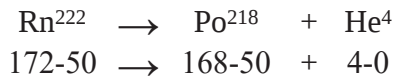
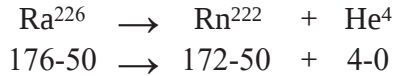
This puts the vibrational mass outside the zone of stability, and two successive beta events follow promptly, bring the atom back to another isotope of uranium.



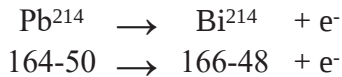
Two successive alpha emissions now take place, with a considerable delay between stages, as both U^{234} and the intermediate product Th^{230} have relatively long half lives. These two events bring the atomic structure to that of radium, the prototype of the radioactive elements.



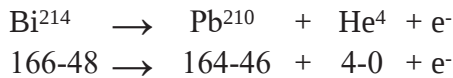
After another somewhat shorter time interval, a rapid succession of decay events begins. Half life periods in this phase of the decay range from days down to as low as seconds. Three more alpha emissions start the sequence.



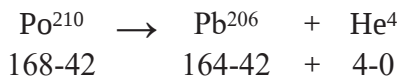
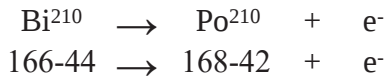
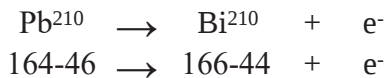
By this time the vibrational mass of 50 units is well above the zone of stability, the center of which is theoretically 43 units at this point. The next event is therefore a beta emission.



This isotope is still above the stable zone, and another beta emission is in order, but a further alpha emission is also imminent, and the next step may take either direction. In either case, the emission is followed by one of the alternate kind, and the net result of the two events is the same regardless of which step is taken first. We may therefore regard this as a double decay.



After some delay due to a 22 year half life of Pb^{210} , two successive beta emissions and one alpha event occur.



The lead isotope Pb^{206} is within the stability limits both with respect to total mass (alpha) and with respect to the ratio of vibrational to rotational mass (beta). The radioactivity therefore ends at this point.

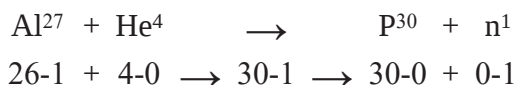
The unstable isotopes that are responsible for natural beta radioactivity in the terrestrial environment originate either as by-products of alpha radioactivity or as a result of atomic transformations originated by high energy processes, such as those initiated by incoming cosmic rays. Alpha radioactivity is mainly the result of past or present inflow of material from regions where the magnetic ionization level is below that of the local environment.

In those regions where the magnetic ionization level is zero, or near zero, all of the 117 possible elements are stable, and there is no alpha radioactivity. The heavy element content of young matter is low because atom building is a cumulative process, and this young matter has not had time to produce more than a relatively small number of the more complex atoms. But probability considerations make it inevitable that *some* atoms of the higher groups will be formed in the younger aggregates, particularly where older matter dispersed into space by explosive processes has been accreted by these younger structures. Thus, although aggregates composed primarily of young matter have a much lower heavy element content than those composed of older matter, they do contain an appreciable number of the very heavy elements, including the trans-uranium elements that are absent from terrestrial matter. The significance of this point will be explained in Volume III.

If matter from a region of zero magnetic ionization is transferred to a region such as the surface of the earth, where the ionization level is unity or above, the stability limit in terms of atomic number drops, and radioactivity is initiated. Whether the material constituents of the earth acquired the unit magnetic ionization level at the time that the earth assumed its present status as a planet, or reached this level at some earlier or later date is not definitely indicated by the information now available. There is some evidence suggesting that this change took place in a considerably earlier era, but in any event the situation with respect to the activity of the elements now undergoing alpha radioactivity is essentially the same. They originated in a region of zero, or near zero, magnetic ionization, and either remained in that region while the magnetic ionization level increased, or in some manner, the nature of which is immaterial in the present connection, were transferred to their present locations, where they have become radioactive for the reasons stated.

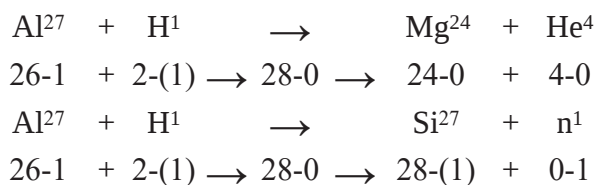
As noted above, another source of natural radioactivity is atomic rearrangement resulting from interaction of material atoms with high energy particles, principally the cosmic rays and their derivatives. In such reactions stable isotopes of one kind or another are converted into unstable isotopes, and the latter then become sources of radioactivity, mostly of the beta type. The level of the beta radioactivity produced in this manner is quite low. The very intense activity of the same general nature that is indicated by the radio and x-ray emission from certain kinds of astronomical objects originates by means of a different process, examination of which will be deferred until the nature and behavior of the objects from which the emissions are observed are developed in Volume III.

The processes that constitute natural radioactivity can be duplicated experimentally, together with a great variety of similar atomic transformations which presumably also occur naturally under appropriate circumstances, but have been observed only under experimental conditions. We may therefore combine our consideration of natural beta radioactivity, the so-called artificial radioactivity, and the other experimentally induced transformations into an examination of atomic transformations in general. Essentially, these transformations, regardless of the number and type of atoms or particles involved, are no different from the simple addition and decay reactions previously discussed. The most convenient way of describing these more complex events is to treat them as successive processes in which the reacting particles first join in an addition reaction and subsequently eject one or more particles from the combination. According to some of the theories currently in vogue, this is the way in which the transformations actually take place, but for present purposes it is immaterial whether or not the symbolic representation conforms to physical reality, and we will leave this question in abeyance. The formation of the isotope P^{30} from aluminum, the first artificial radioactive reaction discovered, may be represented as

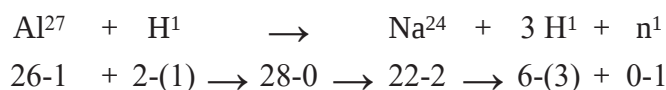


In this case the two phases of the reaction are independent, in the sense that any combination which adds up to 30-1 can produce $P^{30} + n^1$, while there are many ways in which the 30-1 resultant of the combination of $Al^{27} + He^4$ can be broken down. The final product may, for instance, be $Si^{30} + H^1$.

The usual method of conducting transformation experiments is to accelerate a small atomic or sub-atomic unit to a very high velocity and cause it to impinge on a target. In general, the degree of fragmentation of the target atoms depends on the relative stability of those atoms and the kinetic energy of the incident particles. For example, if we use hydrogen atoms against an aluminum target at a relatively low energy level, we get results similar to those produced in the $Al^{27} + He^4$ reaction previously discussed. Typical equations are



Greater energies cause further fragmentation and result in such rearrangements as

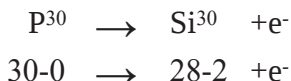


The general principle that the degree of fragmentation is a function of the energy of the incident particles has an important bearing on the relative probabilities of various reactions at very high temperatures, and will have further consideration later.

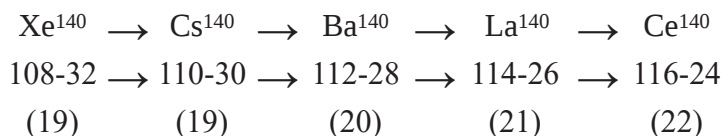
In the extreme situation where the target atom is heavy and inherently unstable, the fragments may be relatively large. In this case, the process is known as *fission*. The difference between the fission process and the transformation reactions previously described is merely a matter of degree, and the same relationships apply.

Although it is possible in some instances to transform one stable isotope into another by an appropriate process, the more general rule is that if the original reactants are stable the major product is

unstable, and therefore radioactive. The reason is, of course, that the stable isotopes have vibrational to rotational mass ratios within the stability zone, and any change in the ratio tends to move it out of that zone. As an example, the P^{30} isotope formed in the reaction between aluminum and helium atoms is below the stability zone; that is, it is deficient in vibrational mass. It therefore decays by the β^+ process to form a stable silicon isotope.



In the radioactive reactions of the heavy elements the products often have substantial excesses of vibrational mass, and in these cases successive beta emissions take place, resulting in decay chains in which the unstable isotopes move step by step toward stability. One of the relatively long chains of this type that has been identified is the following:



The figures in parentheses refer to the number of units of vibrational mass corresponding to the center of the zone of stability, as calculated for each element from equation 24-1. The original product Xe^{140} has 13 excess vibrational units, and is thus far outside the stability zone. Successive beta emissions convert two-unit quantities of vibrational mass to rotational mass, while the stable amount of vibrational mass gradually increases as the atomic number rises. On reaching Ce^{140} the excess has been reduced to two units. This is within the stability margin, and the radioactivity therefore ceases.

The foregoing description of the atomic transformation processes has been confined to the essential element of the transformation, the redistribution of the primary mass, and the collateral effects have either been ignored or left for later treatment. In the latter category are the mass-energy relations, which will be discussed in Chapter 27. The electric charges carried by some of the reacting substances, or the reaction products, have no significance in the present connection, as they only affect the energy relations.

On first consideration it might appear that the addition processes discussed in the preceding pages would provide the answer to the problem of accounting for the existence of the heavier members of the series of chemical elements. In current practice this is taken for granted, and the question to be answered is accepted as being merely the issue as to what specific one or more of these processes is operative.

The currently accepted hypothesis is that the raw material from which the elements are formed is hydrogen, and that mass is added to hydrogen by means of the addition processes. It is recognized that (with certain exceptions that will be considered later) the addition mechanisms are high energy processes. Atoms approaching each other at low or moderate speeds normally rebound, and take up positions at equilibrium distances. The additions take place only where the speeds are high enough to overcome the resistance, and these speeds generally involve disruption of the structure of the target atoms, followed by some recombination.

The only place now known in our galaxy where the energy concentration is at the level required for operation of these processes on a major scale is in the interiors of the stars. The accepted hypothesis therefore is that the atom building takes place in the stellar interiors, and that the products are subsequently scattered into the environment by supernova explosions. It has been demonstrated by laboratory experiments, and more dramatically in the explosion of the hydrogen bomb, that the mass 2 and mass 3 isotopes of hydrogen can be stimulated to combine into the mass 4 isotope of helium, with the release of large quantities of energy. This hydrogen conversion process is currently the most powerful source of energy known to science (aside from some highly speculative ideas that involve carrying gravitational attraction to hypothetical extremes). The attitude of the professional physicists has always been that the most energetic process known to them must necessarily be the process by which energy is generated in the stars (even though they have had to revise their concept of the nature of this process twice already, the last time under very embarrassing circumstances). The current belief of both the physicists and the astronomers therefore is that the hydrogen conversion process is unquestionably the primary stellar

energy source. It is further assumed that there are other addition processes operating in the stars by which atom building beyond the helium level is accomplished.

It will be shown in Volume III that there is a mass of astronomical evidence demonstrating conclusively that this hydrogen conversion process cannot be the means by which the stellar energy is generated. But even without this evidence that demolishes the currently accepted assumption, any critical examination of the fundamentals of atom building will make it clear that high energy processes—inherently destructive—are not the answer to the problem. It is true that the formation of helium from isotopes of hydrogen proceeds in the right direction, but the fact is that the increase in atomic mass that results from the hydrogen conversion reaction is an incidental effect of a process that operates toward a different end. The primary objective of that process, the objective that supplies the probability difference that powers the process, is the conversion of unstable isotopes into stable isotopes.

The fuel for the *known* hydrogen conversion process, that of the hydrogen bomb and the experiments aimed at developing fusion power, is a mixture of these unstable hydrogen isotopes. The operating principle is merely a matter of speeding up the conversion, causing the reactants to do rapidly what they will do slowly, of their own accord, if not subjected to stimulation. It is freely asserted that this is the same process as that by which energy is generated in the stars, and that the fusion experiments are designed to duplicate the stellar conditions. But the hydrogen in the stars is mainly in the form of the stable mass one isotope, and there is no justification for assuming that this stable atomic structure can be induced to undergo the kind of a reaction to which the unstable isotopes are subject by reason of their instability. The mere fact that the conversion process would be exothermic, if it occurred, does not necessarily mean that it will take place spontaneously. The controlling factor is the relative probability, not the energy balance, and so far as we know, the mass one isotope of hydrogen is just as probable a structure as the helium atom under any physical conditions, other than those, to be discussed in Chapter 26, that lead to atom building.

At high temperatures the chances of atomic break-up are improved, but this does not necessarily increase the proportion of helium in the final product. On the contrary, as noted earlier, a greater kinetic energy results in more fragmentation, and it therefore favors the smaller unit rather than the larger. A certain amount of recombination of the fragments produced under these high temperature conditions can be expected, particularly where the extreme conditions are only temporary, as in the explosion of the hydrogen bomb, but the relative amounts of the various possible products of recombination are determined by probability considerations. Inasmuch as stable isotopes are more probable than unstable isotopes (that is what makes them stable), formation of the stable helium isotope from the atomic and sub-atomic fragments takes precedence over recombination of the unstable isotopes of hydrogen. But the mass one hydrogen isotope that is the principal constituent of the stars is just as stable as helium, and it has the advantage, in a high energy environment, of being the smaller unit, which makes it less susceptible to fragmentation, and more capable of recombination if disrupted. Thus it cannot be expected that recombination of fragments into helium, under high energy conditions, will occur on a large enough scale to constitute a major source of stellar energy.

In this connection, it should be noted that the general tendency of high energy reactions in the material sector of the universe is to break down existing structures rather than build larger ones. The reason for this should be evident. The material sector is the low speed sector, and the lower the speed of matter the more pronounced its material character becomes; that is, the more it deviates from the speeds of the cosmic sector. It follows that, in general, the lower the speed the greater the tendency to form combinations of the material type. Conversely, higher speeds lessen the material character of the matter, and not only inhibit further combination, but tend to disrupt the combinations already existing. Furthermore, this increase in the amount of negative displacement (thermal or translational motion) is not conducive to building up positive displacement in the form of mass. Thus the net result of the reactions in the high speed environment of the stellar interiors can be expected to decrease, rather than increase, the average atomic weight of the matter participating in these reactions.

An analogous process in a more familiar energy range is the pyrolysis of petroleum. Cracking of a paraffinic oil, for instance, yields products that, among other things, include substantial quantities of complex aromatic compounds. For example, one of those that makes it appearance is anthracene, a 24-atom molecule. There are few, if any, of the ring compounds, even the smaller ones, in the original material. Thus it is evident that the high temperatures of this process have not only broken down the original hydrocarbon molecules into smaller molecules or atoms, but have also allowed some recombination into larger molecular units. Nevertheless, the general result of the cracking process is a drastic reduction in the average size of the molecules, the greater part of the mass being reduced to hydrogen, methane, and carbon.

The point that needs to be recognized is that this is what high energy processes do to combinations such as atoms, regardless of whether those atoms are combinations of particles, as contended by conventional physics, or combinations of different forms of motion, as deduced from the postulates of the theory of the universe of motion. Such processes disrupt some or all of the original combinations. In the chaotic conditions generated by the application of powerful forces there is a certain amount of recombination going on alongside the disintegration. This may result in the appearance of some new combinations (isotopes), which may suggest that atom building is occurring. But, in fact, these constructive events are merely incidental results of a destructive process.

In the universe of motion, the raw material for atom building consists of massless particles, the decay products of the cosmic rays. Conversion of these particles into simple atoms of matter, and production of increasingly more massive atoms from the original units, is a slow and gradual constructive process, not a high energy destructive process. This assertion as to the general character of the atom building process is confirmed by the astronomical evidence, which, as will be brought out in Volume III, shows that atom building is taking place throughout the universe, not merely at special locations and under special conditions, as envisioned in present-day theories. The details of the atom building processes in the universe of motion will be the subject of the next chapter.

CHAPTER 26

Atom Building

Several chapters of Volume I were devoted to tracing the path followed by the matter that is ejected into the material sector of the universe from the inverse, or cosmic, sector in the form of cosmic rays. As brought out there, the cosmic atoms that constitute the cosmic rays, three-dimensional rotational combinations with net speeds greater than unity, are broken down into massless particles; that is, particles with effective rotation in less than three dimensions. These particles are then reassembled into material atoms, three-dimensional rotational combinations with net speeds less than unity. The processes by which this rebuilding is accomplished have not yet been observed, nor has the applicable theory been fully clarified. It was stated in the earlier volume that our conclusions in this area were necessarily somewhat speculative. Additional theoretical development in the meantime has placed these conclusions on a much firmer basis, and it would now be in order to call them tentative rather than speculative.

As brought out in Chapter 25, the currently prevailing opinion is that atom building is carried on by means of addition processes of the type discussed in that chapter. For the reasons that were specified, we find it necessary to reject that conclusion, and to characterize these processes, to the extent that they actually occur, as minor and incidental activities that have no significant influence on the general evolutionary pattern in the material sector of the universe. However, as noted in the earlier discussion, there is one addition process that actually does occur on a large enough scale to justify giving it some consideration before we turn our attention to broadening the scope of the explanation of the atom building process introduced in Volume I. This addition process that we will now want to examine is what is known as “neutron capture.”

The observed particle known as the “neutron” is the one that we have identified as the compound neutron. It has the same type of structure as the mass one hydrogen isotope; that is, it is a double rotating system with a proton type rotation in one component and a neutrino type rotation in the other. In the hydrogen isotope the neutrino rotation has the material composition $M_{1/2-1/2}-(1)$. In the compound neutron it has the cosmic composition $C(1/2)-(1/2)-1$. The net displacements of this particle are $M_{1/2-1/2}-0$, the same as those of the massless neutron.

The compound neutron is fully compatible with the basic magnetic (two-dimensional) rotational displacement of the atoms, and since it carries no electric charge it can penetrate to the vicinity of an atom much more easily than the particles that normally interact in the charged condition. Consequently, the compound neutrons are readily absorbed by atoms. On first consideration, therefore, neutron capture would appear to be a likely candidate for designation as the primary atom building process. Nevertheless, the physicists relegate it to a minor role. The prevailing downgrading of the potential of neutron capture is mainly due to the physicists’ commitment to other processes that they believe to be responsible for the energy production in the stars. If, as now believed, the continuing additions to the atomic masses are made as a collateral feature of the stellar energy production processes, neutron capture can have only a limited significance. Some support for this conclusion is derived from the finding that there is no stable isotope of mass 5. As the textbooks point out, the neutron capture process would come to a stop at this point.

In the universe of motion this argument is invalid. As we saw in Chapter 24, isotopic stability is determined by the level of magnetic ionization. The lack of a stable isotope of mass 5 is peculiar to the unit ionization level, the level that happens to exist at the surface of the earth at the present time. In earlier eras, when the magnetic ionization level was lower, the obstacle at mass 5 was absent, or at least not fully effective, and in the future when the ionization level has risen, this obstacle will again be minimized or removed.

We must nevertheless concur with the prevailing opinion that neutron capture is not the primary atom building process, because even though the mass 5 obstacle can be circumvented, there are not anywhere near enough of the compound neutrons to take care of the

atom building requirements. These particles are produced in limited quantities in reactions of a special nature. Atom building, on the other hand, is an activity of vast proportions that is going on continuously in all parts of the universe. The compound neutron is actually a very special kind of combination of motions. The reason for its existence is that there are some physical circumstances under which two-dimensional rotation is ejected from matter. In the material atoms the two-dimensional rotation is associated with mass because of the way in which it is incorporated into the atomic structure. There is no way in which this mass can be given up, because the process by which it originated, bringing a massless particle to rest in the fixed spatial reference system, is irreversible. The two-dimensional speed displacement is therefore forced into the only available alternative, the compound neutron structure, even though this structure is inherently one of low probability.

Let us turn now to the process which, according to the findings reported in Volume I, is, in fact, the primary means whereby atom building is actually accomplished. As brought out in that earlier discussion, the principal product of the decay of cosmic atoms, the original constituents of the cosmic rays, is the massless neutron, $M^{1/2-1/2-0}$. This particle can combine with an electron, $M^{0-0-(1)}$, or eject a positron, M^{0-0-1} , to form a neutrino, $M^{1/2-1/2-(1)}$. On the basis of the principles governing the combination of motions, as defined in Volume I, simple combinations of motions do not produce stable structures unless the added motion has some characteristic opposed to that of the original. However, this restriction does not apply to a combination with a neutrino, as this particle has a net total speed displacement of zero, and the added motion is therefore the only active unit in the combination. Thus a massless neutron can be added to a neutrino. Some significant consequences ensue.

All massless particles are moving outward at the speed of light (unit speed) relative to the conventional spatial reference system. But when the neutrino, $M^{1/2-1/2-(1)}$, combines with the massless neutron, $M^{1/2-1/2-0}$, the displacements of the combination are $M^{1-1-(1)}$, which means that the combination has an active inward two-dimensional rotational displacement in a three-dimensional type of structure. The addition of inward motion in the third scalar dimension brings

the consolidated particle to rest in the spatial reference system. The results of this sequence of events were described in Volume I. As noted there, although the massless neutron and the neutrino have no effective mass, they do have the two-dimensional analog, t^2/s^2 , of the three-dimensional property, t^3/s^3 , that is known as mass. When one of these particles, moving at the speed of light relative to the spatial reference system comes to rest in the gravitationally bound system represented by the reference coordinates, the unit translational speed thereby eliminated provides the necessary energy, $1/s$, to convert the two-dimensional quantity, the internal momentum, as we have called it, to the three-dimensional quantity, mass.

The product of this process, with rotational displacements 1-1-(1) and a mass of one atomic weight unit, is the *proton*. In conventional physics the proton is regarded as a positively* charged particle that constitutes the nucleus of the hydrogen atom. We find that it is, in fact, a particle, which may or may not carry a positive* electric charge. We also find that as *a particular kind of motion* (not as a particle) it is a constituent of the hydrogen atom. It is not, however, a “nucleus.” The mass one hydrogen isotope is a double rotating system in which the proton type of motion is combined with a motion of the neutrino type. The atom is formed by direct combination of the proton and the neutrino, but the existence of the particles as particles terminates when the combination takes place. At this point the motions that previously constituted the particles become constituent motions of the combination structure, the atom.

This is an appropriate point at which to make some general comments about the successive combinations of different types of motions that are the essence of the atom building process. The key to an understanding of this situation is a recognition of the fact that these are *scalar* motions. The only inherent property of a scalar motion is its positive or negative magnitude, and the representation of that magnitude in the spatial reference system is subject to change in accordance with the conditions prevailing in the environment. The *same* scalar motion can be either translational, rotational, vibrational, or a rotational vibration, and it is free to switch from one of these to another to conform to changed conditions. Such a change is a zero energy process, as previously defined, merely a rearrangement.

This is the same kind of a situation that we encountered in Chapter 17 in connection with ionization. As noted there, ionization of a particle can take place by means of any one of a number of different processes—*absorption* of radiant energy, *capture* of electrons, *contact* with fast moving particles, etc. Since the motions that are involved are of different types, it might appear that we are confronted with a difficult problem when we attempt to explain these processes as interchange of motions. But the situation is simple when it is viewed in scalar terms. The only *inherent* property of these scalar motions—the vibratory photon motion, the rotational electron motion, the translational motion of the atom or particle—is the magnitude. It follows that the magnitude is the only property that is necessarily transmitted unchanged in an interaction. The coupling to the reference system that distinguishes the photon from the electron, or from translational motion, is free to conform to the new environment. In ionization it takes the form of a rotational vibration, regardless of the type of the antecedent motion.

Production of the hydrogen atom in the manner described in the preceding pages terminates the role of the direct addition processes in atom building. The essential step in this process is to bring the massless neutrons from their normal motion at the speed of light (stationary in the natural reference system) to a condition of rest in the fixed spatial reference system. As pointed out in Volume I, this requires the existence of rotational motion in all three scalar dimensions, since the particle is capable of moving at the speed of light (relative to the spatial reference system) in any vacant dimension. The massless neutron does not have the necessary three dimensions of motion, but combination with the neutrino provides the required addition to the neutron dimensions. This combination, 1-1-(1), has a net total three-dimensional rotational displacement (mass) of one unit.

The 1-1-(1) particle, the proton, thus produced cannot accept another massless neutron because of the two-dimensional nature of that particle. Nor can it accept a combination of the massless neutron with a neutrino, as that combination constitutes another proton, and consolidation of two protons is subject to the opposing factors previously considered in connection with the direct combination of atoms. Beyond the mass one hydrogen stage, therefore, atom

building takes place mainly by means of an ionization process that we will now consider.

The neutrinos in the decay products of the cosmic rays are subject to contacts with other particles, particularly photons of radiation. Some of these contacts result in magnetic ionization; that is, a two-dimensional rotational vibration is imparted to the neutrino. Since this is a one-unit displacement in opposition to the one unit of two-dimensional rotational displacement in the neutrino, the resultant net rotational displacement in these two dimensions is zero. As can readily be seen, such a charge could not be applied to a massless neutron. This particle already has zero displacement in the electric dimension, and if the one unit in the magnetic dimensions is neutralized the particle would have no effective speed displacement, and would be reduced to the status of the rotational base, the rotational equivalent of nothing at all. At the primitive level magnetic ionization is therefore confined to the neutrino.

The magnetic ionization process was discussed at length in Chapters 24 and 25, and the steps through which the original ionization of the neutrinos is passed on to the atoms were described in considerable detail. At this time we will take a look at the mass relations, with the objective of demonstrating that the process by which mass is added during the events previously described is irreversible (up to the destructive limits defined in Chapter 25), and that magnetic ionization is therefore an atom-building process of such broad scope that it is clearly the predominant means of accomplishing the formation of the heavier elements.

As explained previously, since the magnetically charged neutrino has no active speed displacement other than the one negative unit in the electric dimension, it is, in effect, a rotating unit of space vibrating in the magnetic dimensions. A material atom, which is a time structure (net displacement in time), can exist in this space of the neutrino just as in any other space. Such an atom is continually moving from one space unit to another. If it enters the space of a neutrino, the rotational vibration of the space unit (the neutrino) is equivalent to, and in equilibrium with, a similar, but oppositely directed, rotational vibration of the atom. When the atom again passes into another space

unit it is a matter of chance whether the vibration goes with it, or is left with the space unit (the neutrino). Thus some of the magnetic charges originally imparted to the neutrinos in a material aggregate are transferred from the neutrinos to the atoms.

Neutrinos, whether charged or uncharged, move at unit speed relative to the spatial reference system, and their occasional periods of coincidence with atoms of matter are possible only because of the finite magnitude of the units of space and time. If the magnetic charge stays with the atom when the atom and neutrino separate, the charge, which is moving at unit speed while it is associated with the neutrino, is brought to rest in the spatial reference system. Elimination of the unit of outward speed is equivalent to adding a unit of inward speed (derived from the kinetic energy of the atom). This neutralizes the outward vibrational unit and leaves only a unit of magnetic (two-dimensional) speed displacement to be absorbed by the atom. Inasmuch as this unit that is absorbed has only half the mass of the full rotational unit, and has no rotation at all in the third dimension, it enters the atom as a unit of vibrational mass. If this puts the isotopic weight of the atom outside the zone of stability, some of the vibrational mass is converted to rotational mass in the manner previously described, moving the atom to a position higher in the atomic series.

The transition from the massless state (stationary in the natural reference system) to the material status cannot be reversed in the material environment, as there is no available process for going directly from rotation to translation. The sub-atomic particles are subject to neutralization reactions in which oppositely directed rotations cancel each other, causing their speed displacements to revert to the translational status. But direct combination of two multi-unit atoms is difficult to accomplish. Because of the reversed direction of the forces in the time region, there is a strong force of repulsion between two such structures when they approach each other. Furthermore, each atom is a combination of motions in different scalar dimensions, and even if two atoms acquire sufficient relative speed to overcome the resistance and make effective contact, they cannot join unless the displacements in the different dimensions reach the proper conditions for combination simultaneously. With

few, if any, exceptions, the additions to the masses of the atoms are therefore permanent (up to the time that one of the destructive limits is reached).

Here, then, the first application of this atom building process is complete. By means of the successive steps that have been identified, the magnetic rotational speed displacement of the massless neutron produced by cosmic ray decay (the only active property of that particle) is converted into an addition to the mass of an atom. Successive additions of the same kind move the atom up the atomic series.

Atom building in intergalactic space is slow because of the low density of matter, but the amount of time spent in this stage is so long that there is sufficient opportunity for production of a finite quantity of all of the 117 possible elements, in proportions determined by the relative probabilities. After this initial period, the existing matter is increasingly concentrated into large aggregates. This speeds up the atom building, but meanwhile there are processes in operation that destroy some of the heavier elements.

A significant aspect of the theoretical findings reported in this and the immediately preceding chapters is the important role of the massless particles, entities which, with the exception of the photon and the neutrino, are not recognized by conventional science. As brought out in the discussion earlier in this chapter, the characteristic feature of these particles is that they have no capability of independent motion, and are therefore stationary in the natural system of reference. It follows that they are moving at unit speed (the speed of light) in the context of the conventional spatial reference system.

According to our findings, there are three categories of material particles (combinations of motions without enough rotational displacement to form the atomic type of structure). These are (1) massless particles, (2) similar particles that have acquired mass, and (3) particles with structures intermediate between those of class (2) and the full atomic structure. Table 36 lists the sub-atomic particles of the material sector.

The mass one hydrogen isotope is included in this list because of its intermediate type structure, although it is generally regarded as a full scale atom. Electric charges that may be present are not shown,

except in the case of the one-dimensional charged particles, where they provide the rotational vibration that brings these particles into the gravitationally bound system. Charges applied to other particles in the list have no significant effect on the phenomena now being considered.

TABLE 36: *THE SUB-ATOMIC PARTICLES*

Massless Particles	
	photon
M 0-0-0	rotational base
M 0-0-(1)	electron
*M $\frac{1}{2}$ - $\frac{1}{2}$ -(1)	charged neutrino
M 0-0-1	positron
M $\frac{1}{2}$ - $\frac{1}{2}$ -(1)	neutrino
M $\frac{1}{2}$ - $\frac{1}{2}$ -0	massless neutron
Particles With Mass	
-M 0-0-(1)	charged electron
+M 0-0-1	charged positron
M 1-1-(1)	proton
Intermediate Systems	
M 1-1-(1)	compound neutron
C ($\frac{1}{2}$)-($\frac{1}{2}$)-1	
M 1-1-(1)	mass 1 hydrogen
M $\frac{1}{2}$ - $\frac{1}{2}$ -(1)	

- * gravitational charge
- negative* electric charge
- + positive* electric charge

An exact duplicate of Table 36 list exists in the cosmic sector, with the speed displacements inverted. In this case the particles are built on the cosmic rotational base, represented as C 0-0-0, rather than the material rotational base, M 0-0-0. The particles not listed in Table 35

that the physicists claim to have discovered—mesons, etc.—are combinations of the cosmic type, either particles from the cosmic sub-atomic list, or full-sized cosmic atoms (where the presumed discoveries are authentic). It is even possible that some of the events of extremely short duration attributed to transient particles may be originated by cosmic chemical compounds.

Recognition of the place of the massless particles in the evolutionary pattern of matter is one of the advances in understanding that has given us the present consistent, and apparently correct, explanation of the transition from cosmic to material (and vice versa). The 1959 publication identified the cyclic nature of the universe, and gave an account of the manner in which the transitions between sectors take place. At that time, however, the existence of the massless particles had not yet been discovered theoretically, and the particle now identified as the compound neutron was thought to be the intermediary by means of which intersector transfer is accomplished. When it was finally realized that the theory requires the existence of a massless neutron, the door to a new understanding of the transition process was opened. It then became evident that the transition is not directly from cosmic to material, but from cosmic (moving inward in time) to neutral (no motion relative to the natural reference system), and then to material (moving inward in space).

This finding revolutionized our concept of the position of the massless particles in the physical picture. It can now be seen that these particles—the neutrino (known to conventional science), the massless electron and massless positron (previously identified as the moving particles in electric currents), the massless neutron, the rotational base, and the gravitationally charged neutrino (discovered theoretically)—are the constituents of a hitherto unknown subdivision of physical existence, a neutral state of the basic units of matter, intermediate between the states of the cosmic and material sectors.

Inasmuch as the atom building process operates by means of successive additions of single units, the relative proportions of the various elements in a material aggregate are directly related to the age of the matter, and inversely related to the atomic number. However, there are a number of collateral factors that modify the

basic relations. As we have seen, production of the mass one isotope of hydrogen is a relatively simple matter, involving nothing more than a union of two simple particles. The next step is more difficult because it requires the formation of a double system in which there are effective rotational displacements in both components. The great majority of the material atoms are therefore still in the hydrogen stage. The first full double system, helium, atomic number 2, is in second place, as would be expected. Beyond this level, the atomic rotation becomes more complex, and factors other than the required number of additions of mass units introduce numerous irregularities into what would otherwise be a regular decrease of abundance with atomic number.

Evidently a single addition to the atomic rotation introduces a degree of asymmetry that decreases stability, as the even-numbered elements are generally more abundant than the odd-numbered ones. For instance, the ten most abundant elements beyond hydrogen in the earth's crust include seven even-numbered elements, and only three with odd atomic numbers. The zone of isotopic stability is likewise wider in the even-numbered than in the odd-numbered elements, as would be expected if they are inherently more stable. Many of the odd-numbered group have only one stable isotope, and there are five within the 117 element range of the terrestrial environment that have no stable isotope at all (in that environment). On the other hand, no even-numbered element, other than beryllium, has less than two stable isotopes.

The same kind of symmetry effect can be seen in the first additions of rotation in the magnetic dimensions. The positive elements of Group 2A, lithium, beryllium, and boron, are relatively scarce, while the corresponding members of group 2B, sodium, magnesium, and aluminum, are relatively abundant. At higher levels this effect is not apparent, probably because the successive additions to these heavier elements are smaller in proportion to the total mass, while the effects of other factors become more significant.

One of the features of the rotational patterns of the elements that introduces variations in their susceptibility to the addition of mass, and corresponding variations in the proportions in which the different

elements occur in material aggregates, is the change in the magnetic rotation that takes place at the midpoint of each rotational group. For example, let us again consider the 2B group of elements. The first three of these elements are formed by successive additions of positive electric displacement to the 2-2 magnetic rotation. Silicon, the next element, is produced by a similar addition, and the probability of its formation does not differ materially from that of the three preceding elements. Another such addition, however, would bring the speed displacement to 2-2-5, which is unstable. In order to form the stable equivalent, 3-2-(3), the magnetic displacement must be increased by one unit in one dimension. The probability of accomplishing this result is considerably lower than that of merely adding one electric displacement unit, and the step from silicon to phosphorus is consequently more difficult than the additions immediately preceding. The total amount of silicon in existence therefore builds up to the point where the lower probability of the next addition reaction is offset by the larger number of silicon atoms available to participate in the reaction. As a result, silicon should theoretically be one of the most abundant of the post-helium elements. The same considerations should apply to the elements at the midpoints of the other rotational groups, when due consideration is given to the general decrease in abundance that takes place as the atomic number increases.

As we will see in Volume III, there are reasons to believe that the composition of ordinary matter at the end of the first phase of its existence in the material sector, the dust cloud phase, conforms to these theoretical expectations. However, the abundances of the various elements in the region accessible to direct observation, a region in a later stage of development, give us a different picture. The total heavy element content does increase with the age of the matter. A representative evaluation finds the percentage of elements heavier than helium ranging from 0.3 in the globular clusters, theoretically the youngest stellar aggregates that are observable, to 4.0 in the Population I stars and interstellar dust in the solar neighborhood, theoretically the oldest matter within convenient observational range. These are approximations, of course, but the general trend is clear.

The peaks in the abundance curve that should theoretically exist at the midpoints of the rotational groups also make their appearance

at the appropriate points in the lower groups of elements. The situation with respect to carbon is somewhat uncertain, because the observations are conflicting, but silicon is relatively abundant compared to the neighboring elements, as it theoretically should be, and iron, the predominant member of the trio of elements at the midpoint of Group 3A is almost as abundant as silicon. But when we turn to the corresponding members of the 3B group, ruthenium, rhodium, and palladium, we find a totally different situation. Instead of being relatively abundant, as would be expected from their positions in the atomic series just ahead of another increase in the magnetic displacement, these elements are rare. This does not necessarily mean that the relative probability effect due to the magnetic displacement step is absent, as all of the neighboring elements are likewise rare. In fact, *all* elements beyond the iron-nickel group exist only in comparatively minute quantities. Estimates indicate that the combined amount of all of these elements in existence is less than one percent of the existing amount of iron.

It does not appear possible to explain the relative abundances in terms of the probability concept alone. A fairly substantial decrease in abundance compared to iron would be in order if the age of the local system were such as to put the peak of probability somewhere in the vicinity of iron, but this should still leave the ruthenium group among the relatively common elements. The nearly complete elimination of the heavy elements, including this group which should theoretically be quite plentiful requires the existence of some additional factor: either (1) an almost insurmountable obstacle to the formation of elements beyond the iron group, or (2) a process that destroys these elements after they are produced.

There is no indication of the existence of any serious obstacle that interferes with the formation of the heavy elements. So far as we can determine, the atom building process is just as applicable to the heavy elements as to the light ones. The building of the heavy elements is endothermic, but this should not be a serious obstacle, and in any event it does not apply below Group 4A, and therefore has no bearing on the scarcity of the 3B and lower division 3A elements. The peculiar distribution of abundances therefore seems to require the existence of a destructive process that prevents the accumulation

of any substantial quantities of the elements heavier than the iron group, even though they are produced in the normal amounts. We have already seen, in Chapter 17, that such a process exists. This process will be examined in detail in Volume III, where it will be shown that the theoretical results of the process are in full agreement with the observed distribution of abundances of the elements.

The entire atom building process described in this chapter is duplicated in the cosmic sector, with space and time interchanged. Here *inverse mass* is added to move the elements up the cosmic atomic series.

CHAPTER 27

Mass and Energy

The discovery of the mass-energy relation $E = mc^2$ by Einstein was a significant advance in physical theory, and has already had some far-reaching physical applications. It is, of course, entirely consistent with the Reciprocal System of theory. Indeed, this theory provides the explanation of the relation that has heretofore been lacking. It is not always recognized that, in the light of current physical thought, this is a very strange relation. Why should the relation between mass and energy be expressible in terms of speed? Einstein supplied no explanation. He derived the relation from the mathematical expression of his theory of relativity, but a mathematical derivation does not *explain* anything until an *interpretation* of the mathematics gives that derivation a physical meaning. The information that has been missing is now supplied by the Reciprocal System. In the universe of motion defined by that system of theory, mass and energy are both reciprocal speeds, differing only in dimensions, mass being three-dimensional, while energy is one-dimensional. Unit energy is therefore the product of unit mass and the second power of unit speed, the speed of light.

This finding as to the true significance of the mass-energy relation has an important effect on its applicability. It shows that the current belief that a quantity of energy always has a certain mass associated with it is erroneous. Reciprocal speed can exist either as mass, or as energy, but not both simultaneously. A quantity of mass, three-dimensional scalar motion, is equivalent to a quantity of energy, one-dimensional scalar motion, only when three-dimensional motion is actually transformed into one-dimensional motion, or vice versa. In other words, an existing quantity of mass does not correspond to any existing energy, but to the quantity of energy that *would come into existence* if the mass is actually converted into energy.

For this reason, Einstein's hypothesis of an increase in mass accompanying increased velocity is inconsistent with our findings. The kinetic energy increment could increase the mass only if it were converted to mass by some appropriate process, and in that event it would cease to be kinetic energy; that is, the corresponding velocity would no longer exist. Actually, this hypothesis of Einstein's is inconsistent with his valid concept of the conversion of mass into energy, regardless of the point of view from which the question is approached. Mass cannot be an accompaniment of kinetic energy, a quantity that increases as the energy increases, and also an entity that can be converted into kinetic energy, a quantity that increases as the energy decreases. The two concepts are mutually exclusive.

In the theoretical universe of motion now being described, the mass-energy relation is applicable only to those processes in which mass disappears and energy appears, or vice versa. The most familiar process of this kind is the interchange between mass and energy that takes place as a result of radioactivity, or similar atomic transformations. As we saw in Chapter 25, the primary mass is conserved in these reactions. In the radioactive disintegration $\text{Ra}^{226} \rightarrow \text{Rn}^{222} + \text{He}^4$, for example, the total primary mass of the original radium atom was 226. The primary mass of the residual radon atom, 222, and that of the ejected alpha particle, 4, likewise add up to 226. Thus any mass-energy conversion involved in atomic transformations of this kind is confined to the *secondary* mass.

Current scientific opinion regards this secondary mass component as the mass which, according to accepted theory, is associated with the "binding energy" that holds the hypothetical constituents of the hypothetical atomic nucleus together. It must be conceded that this "binding energy" concept fits in very well with the prevailing ideas as to the nature of the atomic structure, but it should be remembered that the entire nuclear concept of the atom is purely hypothetical. No part of it has been verified empirically. Even Rutherford's original conclusion that *most* of the mass of the atom is concentrated in a small nucleus—the hypothesis from which the present-day atomic theory was derived—is not supported by the experiments, except on the basis of the *assumption* that the atoms are in contact in the solid state, an assumption that we now find is erroneous. And every

additional step that has been taken in the long series of adjustments and modifications to which the theory has been subjected as a means of extricating it from difficulties has involved one or more further assumptions, as pointed out in Chapter 18. Thus the fact that the “binding energy” concept is consistent with this aggregate of hypotheses has no physical significance. All available evidence is consistent with our finding that the difference between the observed total mass and the primary mass is a secondary mass effect due to motion within the time region, and that the conversion of this secondary mass to energy is responsible for the energy production during radioactivity or other atomic transformations.

The nature of the secondary mass was explained in Volume I. The magnitudes of this quantity applicable to the sub-atomic particles and the hydrogen isotopes were also calculated. Some studies were made on the higher elements during the early stages of the investigation, and it was shown in the first edition of this work that there is a fairly regular decrease in the secondary mass of the most abundant isotope of the elements in the range from lithium to iron. Beyond iron the values are more irregular, but the secondary mass (negative in this range) remains in the neighborhood of the iron value up to about the midpoint of the atomic series, after which it gradually decreases, and returns to positive values in the very heavy elements. The effect of this secondary mass pattern is to make both the growth process in the light elements and the decay process in the heavy elements exothermic.

From the foregoing, it follows that the secondary mass in the lower half of the atomic series, with the exception of hydrogen, is negative. This conflicts with the general belief that mass is always positive, but our previous development of the theory has shown that the observed mass of an atom is the algebraic sum of the mass equivalents of the speed displacements of the constituent rotations. Where a rotation is negative, the corresponding mass component is also negative. The net total mass of a material atom is always positive only because the magnetic rotation is necessarily positive in the material sector of the universe, and the magnetic rotation is the principal component of the total. Just why the minimum in the secondary mass is at or near the midpoint of the atomic series rather than at one of the extremes is

still unknown, but a similar pattern was noted in some of the material properties examined in the preceding pages of this and the earlier volume, and it is not unlikely that there is a common cause.

Many investigators have devoted considerable effort to the study and analysis of atomic transformations that might possibly serve as the source of the energy generated in the sun and other stars. The general conclusion has been that the most likely reactions are those in which hydrogen is converted into helium, either directly or through a series of intermediate reactions. Hydrogen is the most abundant element in the stars, and in the universe as a whole. This hydrogen conversion process, if actually in operation, could therefore furnish a substantial supply of energy. However, as brought out in Chapter 25, there is no actual evidence that the conversion of ordinary hydrogen, the H^1 isotope, to helium is a *naturally* occurring process in the stars or anywhere else. Even without the new information supplied by the investigation here being reported, there are many reasons to doubt that this process is actually operative, and to question whether it would supply enough energy to meet the stellar requirements if it were in operation. But the energy supply is a critical issue, and since no alternative comes that comes anywhere near meeting the requirements has heretofore been available, the astronomers have evidently felt that they would have to accept the hydrogen process and interpret their own observations to agree with the consequences of that process, even though in many instances this strains credulity to the utmost. As Bart J. Bok says with respect to the conclusions as to the age of the “most conspicuous supergiant” stars that necessarily follows from this hypothesis, “it is no small matter to accept as proven” such a conclusion.¹⁰⁷ The hypothesis obviously fails by a wide margin to account for the enormous energy output of the quasars and other compact astronomical objects. As one astronomer states the case, the problem of accounting for the energy of the quasars “is widely considered to be the most important unsolved problem in theoretical astrophysics.”¹⁰⁸ It will not be feasible to present the astronomical evidence here, but it has had some consideration in previous publications, particularly *New Light on Space and Time*, and the first edition of this work, and it will be covered in detail in Volume III. The point to be noted at this time is that the astronomical

observations give powerful support to the conclusions reached in the previous paragraphs from an analysis of the physical situation.

The catastrophic effect that the invalidation of the hydrogen conversion process as the stellar energy source would otherwise have on astronomical theory, leaving it without any explanation of the manner in which this energy is generated, is avoided by the fact that the development of the Reciprocal System of theory has revealed the existence of not only one, but two hitherto unknown physical phenomena, each of which is far more powerful than the hydrogen conversion process. These newly discovered processes are not only capable of meeting the energy requirements of the stable stars, but also the far greater requirements of the supernovae and the quasars (when the quasar energies are scaled down to the true magnitudes from the inflated values based on the current interpretations of the redshifts of these objects).

Perhaps some readers may find it difficult to accept the thought that there could be hitherto unknown processes in operation in the universe that are vastly more powerful than any previously known process. It might seem that anything of that magnitude should have made itself known to observation long ago. The explanation is that the *results* of these processes *are known observationally*. Extremely energetic events are prominent features of present-day astronomy. What has not been known heretofore is the *nature* of the processes whereby the enormous energies are generated. This is the information that the theory of the universe of motion is now supplying.

In Chapter 17 we examined one of these processes, the conversion of mass to energy that results when the matter in the interior of a star reaches the destructive thermal limit. This is the long-continuing process that supplies the relatively modest (on the astronomical scale) amount of energy necessary to meet the requirements of the stable stars. It also accounts for the large energy output of one kind of supernova, as we will see in Volume III. At this time we will take a look at what happens when a star arrives at a different kind of a destructive limit.

The destructive limit identified in Chapter 17 is reached when the total of the outward displacements (thermal and electric ionization)

reaches equality with one of the inward rotational displacements of the atom, reducing the net displacement of the combination to zero, and destroying its rotational character. A similar destructive limit is reached when the inward displacements (rotation and gravitational charge) are built up to a level that, from the rotational standpoint, is the *equivalent* of zero.

This concept of the equivalent of zero is new to science, and may be somewhat confusing, but its nature can be illustrated by consideration of the principle on which the operation of the stroboscope is based. This instrument observes a rotating object in a series of views at regular intervals. If the interval is adjusted to equal the rotation time, the various features of the rotating object occupy the same positions in each view, and the object therefore appears to be stationary. A similar effect was seen in the early movies, where the wheels of moving vehicles often appeared to stop rotating, or to rotate backward.

In the physical situation, if a rotating combination completes its cycle in a unit of time, each of the displacement units of the combination returns to the same circumferential position at the end of each cycle. From the standpoint of the macroscopic behavior of the motion, the positions at the ends of the time units are the only ones that have any significance—that is, what happens *within* a unit has no effect on other units—and, under the conditions specified, these positions lie in a straight line in the reference system. This means that there is no longer any factor tending to keep the units together as a rotational combination (an atom). Consequently, they separate as linear motions, and mass is transformed into energy. It should be understood, however, that this transformation at the destructive limit has no effect on the motion itself. Scalar motion has no property other than its positive or negative magnitude, and that remains unchanged. What is altered is the coupling to the reference system, which is subject to change at the end of any unit, if the conditions existing at that point are favorable for such a change.

The emphasis on the *ends* of the units of motion in the foregoing discussion is a reflection of the nature of the basic motions, as defined in the fundamental postulates of the Reciprocal System of theory. According to these postulates, the basic units of motion are discrete.

This does not mean that the motion proceeds by a succession of jumps. On the contrary, motion is inherently a continuous progression. A new unit of the progression begins at the point where the preceding unit ends, so that continuity, in this sense, is maintained from unit to unit, as well as within units. But since the units are separate entities, the effects of the events that take place in one unit cannot be carried forward to the next (although the combination of the internal and external features of the *same* unit may be effective, as in the case of the primary and secondary mass). The individual units of motion *may* continue on the same basis, but the coupling of the motion to the reference system is subject to change to conform to whatever conditions may exist at the end of a unit. When the atom has returned to the situation that existed at the original zero, as is true if the end of the rotational cycle coincides with the end of the time unit, the motion has reached a new starting point, a new zero, we may say.

For the reasons previously given, the limiting value, the equivalent of zero in each scalar dimension, is eight units of one-dimensional, or four units of two-dimensional, rotational displacement. In the notation used herein, the latter is a 4-4 magnetic combination. However, as indicated in Chapter 24, the destructive limit is not reached until the displacement in the electric dimension also arrives at the equivalent of the last magnetic unit. A rotational combination (atom) is therefore stable, at zero magnetic ionization, up to 4-4-31, or the equivalent 5-4-(1), which is element 117. One more step reaches the limit at which the rotational motion terminates.

If the rotational limit is reached in atoms whose individual magnetic ionization is above the general level in the aggregate of which these atoms are constituents, the effect of approaching the limit is that the atoms become radioactive, and eject portions of their masses in the form of alpha particles, or other fragments. This prevents the building of elements heavier than number 117, but it does not result in destruction of primary mass such as that which occurs at the destructive thermal limit. Thus the radioactivity is a means of avoiding the destructive effects of reaching the limiting value of the magnetic displacement.

This situation is analogous to a number of others that are more familiar. For example, we saw in Chapter 5 that the limiting value of

the specific heat of a solid is reached at a relatively low temperature. Beyond this limit the atom, or molecule, enters the liquid state. The transition requires a substantial energy input, and since the lower energy states are more probable in a low energy environment, the atom avoids the need to provide the energy increment by changing to a different thermal vibration pattern, if it has the capability of so doing. The atoms of the heavier elements make several changes of this kind as new limiting values of the specific heat are encountered at successively higher temperatures. Eventually, however, a point is reached at which no further expedients of this kind are available, and the atom must pass into the liquid state. Similarly, the probabilities favor the continued existence of the combination of motions that constitutes the atom, as long as this is possible. The destructive effects of arriving at the displacement limit are therefore avoided by the ejection of mass. But here, too, as in the case of the specific heat, a point is eventually reached where the level of magnetic ionization tending to increase the atomic mass prevents further ejection of mass from the atom, and arrival at the destructive limit can no longer be avoided.

The consequences of reaching this rotational displacement limit at the equivalent of zero are qualitatively identical with those of reaching the thermal displacement limit at zero. The various rotational components cancel out, and the motion reverts to the linear basis. This transforms mass into kinetic energy, most of which is imparted to the residue of the atoms, or to other matter in the environment. The remainder goes into electromagnetic radiation. From a quantitative standpoint, there are some significant differences between the two phenomena. The thermal limit applies only to the heaviest element that is present in the aggregate in a significant quantity, and the rate at which this element arrives at the limit is regulated by a process that will be discussed in Volume III. The elements lower in the atomic series are not affected. Furthermore, the conversion of rotational to linear displacement (mass to energy) at the thermal limit does not necessarily apply to more than one of the magnetic displacement units of the atom, and a large part of the atomic mass may therefore remain intact, either as a residual atom or a number of fragments.

Consequently, the thermal limit has no catastrophic effect until the temperature reaches the destructive limit of an element, iron, that is

present in relatively large quantities. On the other hand, arrival at the magnetic displacement limit affects the entire mass of each atom, and the only portion of the mass of an aggregate that remains intact is that in the outer portions of the aggregate where the magnetic ionization level is lower than in the deeper interior. There is no process that limits the rate of disintegration at this destructive limit. The resulting explosion, known as a Type II supernova, is therefore much more powerful (relative to the mass of the exploding star) than the Type I supernova that occurs at the thermal limit, although its full magnitude is not evident from direct observation, for reasons that will be explained in Volume III.

While the thermal disintegration process is operative in every star, it does not necessarily proceed all the way to destruction of the star. The extent to which the mass of the star, and consequently the temperature, increases depends on its environment. Some stars will accrete enough mass to reach the temperature limit and explode; others will not. But the increase in the magnetic ionization level is a continuing process in all environments, and it necessarily results in arrival at the magnetic destructive limit when sufficient time has elapsed. This limit is thus essentially an *age* limit.

A process related to those that have been described in the foregoing paragraphs is the sequence of events that counterbalances the conversion of three-dimensional motion (mass) into one-dimensional motion (energy) in the stars. The energy that is generated by atomic disintegration leaves the stars in the form of radiation. According to present-day views, this radiation moves outward at the speed of light, and most of it eventually disappears into the depths of space. The theory of the universe of motion gives us a very different picture. It tells us that inasmuch as the photons of radiation have no capability of independent motion relative to the natural datum, they remain stationary in the natural reference system, or move inward at the speed of the emitting object. Each photon therefore eventually encounters, and is absorbed by, an atom of matter. The net result of the generation of stellar energy by atomic disintegration is thus an increase in the thermal energy of other matter. As will be explained in Volume III, the matter of the universe is subject to a continuing process of aggregation under the influence of gravitation.

Consequently, all matter in the material sector, with the added thermal energy, is ultimately absorbed by one of the giant galaxies that are the end products of the aggregation process.

When supernova explosions in the interior of one of these giant galaxies become frequent enough to raise the average particle speed above the unit level, some of the full units of speed thus made available are converted into rotational motion, creating cosmic atoms and particles. This cosmic atom building, which theoretically operates on a very large scale in the galactic interiors, has been observed on a small scale in experiments, the results of which were discussed in Volume I. In the experiments, the high energy conditions are only transient, and the cosmic atoms and particles that are produced from the high level of kinetic energy quickly decay into particles of the material system. Some such decays no doubt also occur in the galactic interiors, but in this case the high energy condition is quasi-permanent, favoring continued existence of the cosmic units until ejection of the quasar takes place. In any event, the production of these rotational combinations has increased the amount of existing cosmic or ordinary matter at the expense of the amount of existing energy, thus reversing the effect of the production of energy by disintegration of atoms of matter.

In concluding this last chapter of a volume dealing with the properties of matter, it will be appropriate to call attention to the significant difference between the role that matter plays in conventional physical theory, and its status in the theory of the universe of motion. The universe of present-day physical science is a universe of matter, one in which the presence of matter is the central fact of physical existence. In this universe of matter, space and time provide the background, or setting, for the action of the universe; that is, according to this view, physical phenomena take place *in* space and *in* time.

As Newton saw them, space and time were permanent and unchanging, independent of each other and of the physical activity taking place in them. Space was assumed to be Euclidean (“flat” in the jargon of present-day mathematical physics), and time was assumed to flow uniformly and unidirectionally. All magnitudes, both of space and of time, were regarded as absolute; that is, not dependent on

the conditions under which they are measured, or on the manner of measurement. A subsequent extension of the theory, designed to account for some observations not covered by the original version, assumed that space is filled with an imponderable fluid, the *ether*, which interacts with physical objects.

Einstein's relativity theories, which have replaced Newton's theory as the generally accepted view of the theoretical physicists, retain Newton's concept of the general nature of space and time. To Einstein these entities constitute a background for the action of the universe, just as they did for Newton. Instead of being a three-dimensional space and a one-dimensional time, independent of each other, as they were for Newton, they are amalgamated into a four-dimensional spacetime in Einstein's system, but they still have exactly the same function; they form the framework, or container, within which physical entities exist and physical events take place. Furthermore, these basic physical entities and phenomena are essentially identical with those that exist in Newton's universe.

It is commonly asserted that Einstein eliminated the ether from physical theory. In fact, however, what he actually did was to eliminate the *name* "ether," and to apply the name "space" to the concept previously called the "ether." Einstein's "space" has the same kind of properties that were formerly assigned to the ether, as he admits in the following statement:

We may say that according to the general theory of relativity space is endowed with physical qualities; in this sense, therefore, there still exists an ether.²⁵

The downfall of Newtonian physics was due to a gradual accumulation of discrepancies between theory and observation, the most critical being the results of the Michelson-Morley experiment and the measurements of the advance of the perihelion of Mercury, neither of which could be explained within the limits of Newton's system. Some modification of that system was obviously necessary. The question, as it stood around the end of the nineteenth century, was what form the revision of Newton's ideas should take.

As brought out in Chapter 13, in order to qualify as "theory," in the full meaning of the term, the treatment of a physical phenomenon

must cover not only its mathematical aspects, but also its physical aspects; that is, it must provide a conceptual understanding of the entities and relations to which the mathematics refer. However, the general tendency in recent years has been to concentrate on the mathematical development and to omit the parallel conceptual development, substituting conceptual interpretations of the individual mathematical results. Richard Feynman describes the present situation in this manner:

Every one of our laws is a purely mathematical statement in rather complex and abstruse mathematics.⁵⁶

In his attack on the problem of revising Newton's theory, Einstein not only adopted this policy of widening the latitude for theory construction by restricting his development to the mathematical aspects of the subject under consideration, and thereby avoiding any conceptual limitations on his basic assumptions, but went a step farther, and loosened the normal mathematical constraints as well. He first introduced a high degree of flexibility into the numerical values by discarding "the idea that co-ordinates must have an immediate metrical meaning,"³⁶ an expression that he defines as the existence of a specific relationship between differences of coordinates and measurable lengths and times. As C. Møller describes this theoretical picture:

In accelerated systems of reference the spatial and temporal coordinates thus lose every physical significance; they simply represent a certain arbitrary, but unambiguous, numbering of physical events.¹⁰⁹

Along with this flexibility of physical measurement, which greatly increased the latitude for making additional assumptions, Einstein introduced a similar flexibility into the geometry of spacetime by assuming that it is distorted or "curved" by the presence of matter. The particular aim of this expedient was to provide a means of dealing with gravitation, a key issue in the general problem. One textbook explains the new view in this manner:

What we call a gravitational field is equivalent to a "warping" of time and space, as if it were a rubbery sort of material that stretched out of shape near heavy bodies.¹¹⁰

The basis for this assertion is an *assumption*, the assumption that, for some unspecified reason, space and matter exert an influence upon each other. “Space acts on matter, telling it how to move. In turn, matter acts on space, telling it how to curve.”¹¹¹ (Misner, Thorne, and Wheeler) But neither Einstein nor his successors have given us any explanation of *how* such interactions are supposed to take place—*how* space “tells” matter, or vice versa. Nor does the theory explain inertia, an aspect of the gravitational situation that has given the theorists considerable trouble. As Abraham Pais sums up this situation:

It must also be said that the origin of inertia is and remains the most obscure subject in the theory of particles and fields.¹¹²

Today there is a tendency to call upon Mach’s principle, which attributes the local behavior of matter to the influence of the total quantity of matter in the universe. Misner, Thorne, and Wheeler say that “Einstein’s theory identifies gravitation as the mechanism by which matter there (the distant stars) influences inertia here.”¹¹³ But, as indicated in the statement by Pais, this explanation is far from being persuasive. It obviously gives us no answer to the question that baffled Newton: How does gravitation originate?. Indeed, there is something incongruous about the acceptance of Mach’s principle by the same scientific community that is so strongly opposed to the concept of action at a distance.

The fact is that neither Newton’s theory nor Einstein’s theory tells us anything about the “mechanism” of gravitation. Both take the existence of mass as something that has to be accepted as a given feature of the universe, and both require that we accept the fact that masses gravitate, without any explanation as to how, or why, this takes place. The only significant difference between the two theories, in this respect, is that Newton’s theory gives us no reason why masses gravitate, whereas Einstein’s theory gives us no reason why masses cause the distortion of space that is asserted to be the reason for gravitation. As Feynman sums up the situation, “There is no model of the theory of gravitation today, other than the mathematical form.”⁵⁶

The concept of a universe of motion now provides a gravitational theory that not only explains the gravitational mechanism, but

also clarifies its background, showing that mass is a necessary consequence of the basic structure of the universe, and does not have to be accepted as unexplainable. This theory is based on a new, and totally different, view of the status of space and time in the physical universe. Both Newton and Einstein saw space and time as the container for the constituents of the universe. In the theory of the universe of motion, on the other hand, space and time are the constituents of the universe, and there is no container. On this basis, the space of the conventional spatio-temporal reference system is just a reference system—nothing more. Thus it cannot be curved or otherwise altered by the presence or action of anything physical. Furthermore, since the coordinates of the reference system are merely representations of existing physical magnitudes, they automatically have the “metrical meaning” that Einstein eliminated from his theory to attain the flexibility without which it could not be fitted to the observations.

The theory of the universe of motion is the first physical theory that actually *explains* the existence of gravitation. It demonstrates that the gravitational motion is a necessary consequence of the properties of space and time, and that the same thing that makes an atom an atom, the rotationally distributed scalar motion, also causes it to gravitate. Additionally, the same motion is responsible for inertia.

Of course, this return to absolute magnitudes and mathematical rigidity invalidates the conceptual interpretations of Einstein's solutions of the problems raised by the observed deviations from the consequences of Newton's theory, and requires finding new answers to these problems. But these answers have emerged easily and naturally during the course of the development of the details of the new theory. In most cases no changes in the existing formulation of the mathematical relations have been required. While Einstein's modification of Newton's theory was almost entirely mathematical, our modification of the Newton-Einstein system is primarily conceptual, because the errors in currently accepted theory are nearly all in the conceptual interpretation of the observations and measurements; that is, in the prevailing understanding of the *meaning* of the mathematical terms and the relations between them.

The changes that the new theory makes in the conceptual aspects of the gravitational situation do not affect any of the valid mathematical results of Einstein's theory. For example, most of the mathematical consequences of the general theory of relativity that have led to its acceptance by the scientific community are derived from one of its postulates, the Principle of Equivalence, which states that gravitation is the equivalent of an accelerated motion. In the theory of the universe of motion, gravitation *is* an accelerated motion. It follows that any conclusion that can *legitimately* be drawn from the Principle of Equivalence, such as the existence of gravitational redshifts, can likewise be derived from the postulates of the theory of the universe of motion in exactly the same form.

The agreement between the two theories that exists in these subsidiary areas, and in the mathematical results, does not extend to the fundamentals of gravitation. Here the theories are far apart. The theoretical development reported in the several volumes of this work shows that the attempt to resolve physical issues by mathematical means—the path that has heretofore been followed in dealing with fundamental physics—precludes any significant conceptual changes in theory, whereas, as our findings have demonstrated, there are major errors in the basic assumptions upon which the mathematical theories have been constructed.

Until comparatively recently it was not feasible to locate and correct these errors, because access to a large amount of factual information is indispensable to such an undertaking, and the available supply of information was simply not adequate. Continued research has overcome this obstacle, and the development of the theory of the universe of motion has now identified the “machinery,” not only of gravitation, but of physical processes in general. We are now able to identify the common denominator of all of the fundamental physical entities, and by defining it, we define the entire structure of the physical universe.

References

1. Weisskopf, V. F., *Lectures in Theoretical Physics*, Vol. III, Britten, J. Downs, and B. Downs, editors, Interscience Publishers, New York, 1961, p. 80.
2. Wyckoff, R. W. G., *Crystal Structures*, and supplements, Interscience Publishers, New York, 1948 and following.
3. All compression data in the range below 100.000 kg/cm² not otherwise identified are from the works of P. W. Bridgman, principally his book, *The Physics of High Pressure*, G. Bell & Sons, London, 1958, and a series of articles in the Proceedings of the American Academy of Arts and Sciences.
4. Kittel, Charles, *Introduction to Solid State Physics*, fifth edition, John Wiley & Sons, New York, 1976.
5. McQueen and Marsh, *Journal of Applied Physics*, July 1960.
6. National Bureau of Standards, *Tables of Normal Probability Functions*, Applied Mathematics Series, No. 23, Washington, DC, 1953.
7. Specific heat data are mainly from Hultgren, et al, *Selected Values of the Thermodynamic Properties of the Elements*, American Society for Metals, Metals Park, OH, 1973, and *Thermophysical Properties of Matter*, Vol. 4, Touloukian and Buyko, editors, IFI Plenum Data Co., New York, 1970, with some supplemental data from original sources.
8. Heisenberg, Werner, *Physics and Philosophy*, Harper & Bros., New York, 1958, p. 189.
9. Heisenberg, Werner, *Philosophic Problems of Nuclear Science*, Pantheon Books, New York, 1952, p. 55.
10. Bridgman, P. W., *Reflections of a Physicist*, Philosophical Library, New York, 1955, p. 186.
11. Weisskopf, V. F., "The Frontiers and Limits of Science," *American Scientist*, July-Aug. 1977, pp. 405-406.
12. Values of the coefficient of expansion are from *Thermophysical Properties of Matter*, *op. cit.*, Vol. 12.
13. Duffin, W. J., *Electricity and Magnetism*, second edition, John Wiley & Sons, New York, 1973, p. 122.
14. *Ibid.*, p. 52.
15. Robinson, F. N. H., in *Encyclopedia Britannica*, fifteenth edition, Vol. 6, p. 543.

16. Stewart, John W., in *McGraw-Hill Encyclopedia of Science and Technology*, Vol. 4, p. 199.
17. Landé, Alfred, "Non-Quantal Foundations of Quantum Theory" in *Philosophy of Science*, Vol. 24, No. 4 (Oct., 1957), p. 310.
18. Meaden, G. T., *Electrical Resistance of Metals*, Plenum Press, New York, 1965, p. 1.
19. *Ibid.*, p. 22.
20. The resistance values are from Meaden, *op. cit.*, supplemented by values from other compilations and original sources.
21. Values of the thermal conductivity are from *Thermophysical Properties of Matter*, *op. cit.*, Vol. I.
22. Davies, Paul, *The Edge of Infinity*, Simon & Schuster, New York, 1981, p. 137.
23. Alfvén, Hannes, *Worlds-Antiworlds*, W. H. Freeman & Co., San Francisco, 1966, p. 92.
24. Einstein, Albert and Infeld, Leopold, *The Evolution of Physics*, Simon & Schuster, New York, 1938, p. 185.
25. Einstein, Albert, *Sidelights on Relativity*, E. P. Dutton & Co., New York, 1922, p. 23.
26. *Ibid.*, p. 19.
27. Einstein and Infeld, *op. cit.*, p. 159.
28. Dicke, Robert H., *American Scientist*, Mar. 1959.
29. Bridgman, p. W., *The Way Things Are*, Harvard University Press, 1959, p. 153.
30. Heisenberg, Werner, *Physics and Philosophy*, *op. cit.*, p. 129.
31. Eddington, Arthur S., *The Nature of the Physical World*, Cambridge University Press, 1933, p. 156.
32. Einstein and Infeld, *op. cit.*, p. 158.
33. *Science News*, Jan. 31, 1970.
34. Bridgman, p. W., *The Way Things Are*, *op. cit.*, p. 151.
35. Von Laue, Max, *Albert Einstein, Philosopher-Scientist*, edited by p. A. Schilpp, *The Library of Living Philosophers*, Evanston, IL, 1949, p. 517.
36. Andrade, E. N. da C., *An Approach to Modern Physics*, G. Bell & Sons, London, 1960, p. 10.
37. Carnap, Rudolf, *Philosophical Foundations of Physics*, Basic Books, New York, 1966, p. 234.
38. Gardner, Martin, *The Ambidextrous Universe*, Charles Scribner's Sons, New York, 1979, p. 200.
39. Einstein, Albert, *Albert Einstein: Philosopher-Scientist*, *op. cit.*, p. 21.

40. Duffin, W. J., *op. cit.*, p. 25.
41. *Ibid.*, p. 281.
42. Gerholm, Tor R., *Physics and Man*, The Bedminster Press, Totowa, NJ, 1967, p. 135.
43. *Ibid.*, pp. 147, 151.
44. Davies, Paul, *Space and Time in the Modern Universe*, Cambridge University Press, 1977, p. 139.
45. Einstein, Albert, *Albert Einstein: Philosopher-Scientist*, *op. cit.*, p. 67.
46. Lorrain, Paul and Corson, Dale R., *Electromagnetism*, W. H. Freeman & Co., San Francisco, 1978, p. 95.
47. Maxwell, J. C., *Royal Society Transactions*, Vol. CLV.
48. Rojansky, Vladimir, *Electromagnetic Fields and Waves*, Dover Publications, New York, 1979, p. 280.
49. Smythe, William R., *McGraw-Hill Encyclopedia*, *op. cit.*, Vol. 4, p. 338.
50. Kip, Arthur, *Fundamentals of Electricity and Magnetism*, second edition, McGraw-Hill Book Co., New York, 1969, p. 136.
51. Duffin, W. J., *op. cit.*, p. 301.
- [51a. The correctness of these dimensions of capacitance in relation to the electric current system has been demonstrated by Ronald W. Satz in a paper entitled "Theory of the Capacitor," included as chapter 2-3e in *Existsents and Interactions: A Computational Treatise of the Reciprocal System* (Transpower Corporation, Pennndel, PA, 2017)]
52. *Ibid.*, p. 313.
53. *Ibid.*, p. 302.
54. Bleaney, B.I. and Bleaney, B., *Electricity and Magnetism*, third edition, Oxford University Press, New York, 1976, p. 64.
55. Dobbs, E. R., *Electricity and Magnetism*, Routledge & Kegan Paul, London, 1984, p. 50.
56. Feynman, Richard, *The Character of Physical Law*, MIT Press, 1967, p. 39.
57. Duffin, W. J., *op. cit.*, p. 3.
58. Rogers, Eric M., *Physics for the Inquiring Mind*, Princeton University Press, 1960, p. 550.
59. McCaig, Malcolm, in *Permanent Magnets and Magnetism*, D. Hadfield, editor, John Wiley & Sons, New York, 1962, p. 18.
60. Park, David, *Contemporary Physics*, Harcourt, Brace & World, New York, 1964, p. 122.
61. Mellor, J. W., *Modern Inorganic Chemistry*, Longmans, Green & Co., London, 1919.
62. Rogers, Eric M., *op. cit.*, p. 407.

63. Park, David, *op. cit.*, p. 15.
64. Margenau, Henry, *Open Vistas*, Yale University Press, 1961, p. 72.
65. Thorne, Kip S., *Scientific American*, Dec. 1974.
66. Misner, Thorne, and Wheeler, *Gravitation*, W. H. Freeman & Co., New York, 1973, p. 620.
67. Feynman, Richard, *The Feynman Lectures on Physics*, Vol. II, Addison-Wesley Publishing Co., Menlo Park, CA, 1977, p. 1-3.
68. Bueche, F. J., *Understanding the World of Physics*, McGraw-Hill Book Co., New York, 1981, p. 328.
69. Pais, Abraham, *Subtle is the Lord*, Oxford University Press, New York, 1982, p. 34.
70. Hanson, Norwood R., in *Scientific Change*, edited by A. C. Crombie, Basic Books, New York, 1963, p. 492.
71. Schrödinger, Erwin, *Science and Humanism*, Cambridge University Press, 1952, p. 22.
72. Schrödinger, Erwin, *Science and the Human Temperament*, W. W. Norton & Co., New York, 1935, p. 154.
73. Heisenberg, Werner, *Physics and Philosophy*, *op. cit.*, p. 129.
74. Feynman, Richard, *The Character of Physical Law*, *op. cit.*, p. 168.
75. Jastrow, Robert, *Red Giants and White Dwarfs*, Harper & Row, New York, 1967, p. 41.
76. Harwit, Martin, *Cosmic Discovery*, Basic Books, New York, 1981, p. 243.
77. Ford, K. W., *Scientific American*, Dec. 1963.
78. Hulsizer and Lazarus, *The New World of Physics*, Addison-Wesley Publishing Co., Menlo Park, CA, 1977, p. 254.
79. Bahcall, J. N., *Astronomical Journal*, May 1971.
80. Duffin, W. J., *op. cit.*, p. 165.
81. Hulsizer and Lazarus, *op. cit.*, p. 255.
82. Mc Caig, Malcolm, *op. cit.*, p. 35.
83. Anderson, J. C., *Magnetism and Magnetic Materials*, Chapman and Hall, London, 1968, p. 1.
84. McCaig, Malcolm, *Permanent Magnets in Theory and Practice*, John Wiley & Sons, New York, 1977, p. 339.
85. Duffin, W. J., *op. cit.*, p. 156.
86. McCaig, Malcolm, *Permanent Magnets*, *op. cit.*, p. 341.
87. *Ibid.*, p. 8.
88. Shortley and Williams, *Elements of Physics*, second edition, Prentice-Hall, New York, 1955, p. 717.

89. Lorrain and Corson, *op. cit.*, p. 360.
90. Feynman, Richard, *The Feynman Lectures on Physics*, *op. cit.*, Vol. II, p. 13-1.
91. Bueche, F. J., *op. cit.*, p. 259.
92. *Ibid.*, p. 267.
93. Rogers, Eric M., *op. cit.*, p. 562.
94. Feather, Norman, *Electricity and Matter*, Edinburgh University Press, 1968, p. 104.
95. Kip, Arthur, *op. cit.*, p. 278.
96. *Ibid.*, p. 283.
97. Duffin, W. J., *op. cit.*, p. 210.
98. Lorrain and Corson, *op. cit.*, p. 334.
99. *Handbook of Chemistry and Physics*, sixty-sixth edition, Chemical Rubber Publishing Co., Cleveland, 1976.
100. Martin, D. H., *Magnetism in Solids*, MIT Press, 1967, p. 9.
101. Dobbs, E. R., *op. cit.*, p. 54.
102. Lorrain and Corson, *op. cit.*, p. 237.
103. Duffin, W. J., *op. cit.*, p. 60.
104. *Concepts in Physics*, CRM Books, Del Mar, CA, 1979, p. 266.
105. Walter R. Fuchs, *Physics for the Modern Mind*, The Macmillan Co., NY, 1967, p. 115.
106. Davies, Paul, *The Accidental Universe*, Cambridge University Press, 1982, p. 60.
107. Bok, Bart J., *The Astronomer's Universe*, Cambridge University Press, 1958, p. 91.
108. Mitton, Simon, *Astronomy & Space*, Vol. I, edited by Patrick Moore, Neale Watson Academic Publishers, New York, 1972, p. 304.
109. Møller, C., *The Theory of Relativity*, Clarendon Press, Oxford, 1952, p. 226.
110. Hulsizer and Lazarus, *op. cit.*, p. 222.
111. Misner, Thorne, and Wheeler, *op. cit.*, p. 5.
112. Pais, Abraham, *op. cit.*, p. 288.
113. Misner, Thorne, and Wheeler, *op. cit.*, p. 543.

BASIC PROPERTIES OF MATTER is the second volume of the revised three-volume edition of Dewey B. Larson's initial presentation of the Reciprocal System of Theory, entitled THE STRUCTURE OF THE PHYSICAL UNIVERSE (North Pacific Publishers, 1959).

The present publication supersedes the preliminary edition of this volume, published in a limited edition in 1988, which was marred by numerous typographical errors and omissions. This new edition is based on a meticulous review of Larson's original manuscript and on his corrigenda to the preliminary edition.